



More on Network Approaches in Historical Chinese Phonology ()

Johann-Mattis List

► To cite this version:

Johann-Mattis List. More on Network Approaches in Historical Chinese Phonology (). LFK Society
Young Scholars Symposium, Jul 2018, Taipei, Taiwan. hal-01706927v2

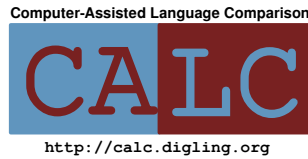
HAL Id: hal-01706927

<https://hal.science/hal-01706927v2>

Submitted on 10 Mar 2018 (v2), last revised 7 Aug 2018 (v3)

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Paper prepared for the *LFK Society Young Scholars Symposium* (Taipei)

More on Network Approaches in Historical Chinese Phonology (音韻學)

Johann-Mattis List^{a*}

February 2018

^aResearch Group “Computer-Assisted Language Comparison”, Department of Linguistic and Cultural Evolution, Max Planck Institute for the Science of Human History

Abstract

The discipline of Historical Chinese Phonology has made great progress during the last thirty years. Thanks to these improvements we have now a much clearer picture of the history of the Chinese language and the development of the Chinese writing system. Since Historical Chinese Phonology is an inherently data-driven discipline, it is, however, surprising that most of the research still heavily relies on qualitative investigations. Given that quantitative approaches have been proven useful to tackle different problems in general historical linguistics, it seems worthwhile to explore how they could be used to help investigating specific problems in Historical Chinese Phonology. Building on recent attempts to apply network approaches to problems of rhyme analysis in Old Chinese poetry and the modeling of Chinese dialect history, this paper presents some new ideas by which specific problems of Historical Chinese Phonology, such as the modeling of character formation processes in the history of the Chinese writing system, the analysis of sound glosses in Middle Chinese literature, or the evaluation of reconstruction systems and rhyme analyses, can be handled with help of network techniques which were originally designed to study problems in other research fields. The ideas presented are understood as work-in-progress. The goal is not to provide full-fledged solutions and applications, but rather to inspire a discussion among colleagues that may help to further improve the methods presented. All examples are accompanied by data and code, enabling scholars interested in the approaches to test the analyses themselves and develop them further.

Keywords

network approaches, Historical Chinese Phonology, Chinese character formation, *fǎnqiè* spellings, reconstruction of Old Chinese phonology, data-display networks

*My research was supported by the DFG research fellowship grant 261553824 and the ERC Starting Grant 715618 CALC. I thank Nathan W. Hill, Laurent Sagart, Guillaume Jacques, and Yunfan Lai for inspiring comments and discussions in the past. This research heavily profited from the intensive collaboration with Nathan W. Hill and Kelsey Granger, with whom I closely collaborated in digitizing and annotating parts of the data which have been used in this paper. I am further grateful to all people who have been working hard on creating datasets and free software libraries and making them publicly available. Without their resources, nothing of what I present in this paper would have been possible.

1 Introduction

Historical Chinese Phonology (*Yīnyùnxué* 音韻學) is a venerable field of research which has successfully elucidated many questions about the development of Chinese dialects, the phonology of ancient stages of Chinese, like Middle and Old Chinese, as well as the development of the Chinese writing system. Especially during the last three decades, great improvements have been achieved, not only in the field of Old Chinese phonology (Baxter 1992, Baxter and Sagart 2014, Starostin 1989, Zhèngzhāng 2003), but also in dialectology (Norman 2003) and the analysis of the Chinese writing system (Qiú 1988 [2007]). The majority of these improvements, however, were based on individual scholarship thanks to which new patterns could be identified. Furthermore, the large part of these improvements was achieved through *qualitative* analysis, although most approaches in Historical Chinese Phonology have always been data-driven in their nature and should thus be amenable to quantitative analysis as well.

Although they have not yet been included in the canon of classical approaches in Historical Chinese Phonology, quantitative approaches have been more frequently applied in recent times, following the more general *quantitative turn* in historical linguistics (Hill and List 2017). Interestingly, quite a few of these approaches are based on network techniques, which were applied to study reticulate evolution in Chinese dialect history (Ben Hamed and Wang 2006, List 2015, List et al. 2014), to detect partially related words in Chinese dialects (List et al. 2016b), or to model rhyming patterns in the Book of Odes (*Shījīng* 詩經) for the purpose of the reconstruction of Old Chinese phonology (List 2017b, List et al. 2017b).

Network approaches are particularly useful for quantitative approaches in historical linguistics and beyond. Mathematically, networks are well understood, and graph theory offers a large arsenal of theoretical frameworks and practical techniques (Newman 2010). Given that many phenomena of our daily lives are also inherently network-like, as reflected in *social networks*, the *world-wide-web*, but also in biological and cultural evolution, it is often easy to model a problem in a particular field as a network problem, in order to search for solutions which have been already proposed in graph theory. It seems therefore like a straightforward question to ask whether networks could be also fruitful to study general problems of Historical Chinese Phonology, which have so far been largely qualitatively investigated.

Given the scope of the paper and the state of the art, these ideas are not to be confused with full-fledged applications and solutions, but rather considered as work-in-progress. Although I offer initial examples for implementations, which are accompanied by software and data illustrating how these ideas could be implemented, it is important to keep in mind that they are only a first step, which will need future refinement and elaboration. On the short run, I hope that I can provide initial proofs of concept for the usefulness of quantitative network approaches to Historical Chinese Phonology. On the long run, I hope that – by illustrating how a quantitative perspective can yield new analyses that might further broaden our knowledge of Historical Chinese Phonology – these ideas will inspire future research, especially by younger colleagues in the field.

It should be clear that all quantitative approaches have certain limits. Pressing complex phenomena into a model will necessarily sacrifice some detail, and it is obvious that what can be captured in pure prose, written by highly educated academics will in its essence always be superior to quantitative models which break down reality to a larger extent. Nevertheless, we should also keep in mind that quantitative approaches offer us a perspective which we could never reach when concentrating on all details of a given problem. This perspective is not limited to the bird's eye perspective – the bigger picture – that may reveal interesting patterns that usually cannot be spotted by eyeballing data with one's own eyes or relying solely on one's intuition. It may also be used to elaborate on a given hypothesis based on qualitative data inspection and substantiate it by providing quantitative backup.

2 Network Approaches to Historical Chinese Phonology

In the following I will try to show how network approaches can help to address couple of problems frequently discussed in Historical Chinese Phonology. These suggestions are drawn from the experience in working with linguistic data in computational frameworks. Researchers who know these techniques may easily recognize that the majority of my suggestions is not new, but instead draws from more general frameworks which have been developed in fields where network approaches are already regularly used to address problems of data modeling and analysis. Not all of the approaches are a success, and I will point to concrete problems and shortcomings when discussing them. All approaches are accompanied by data and source code which are available from the supplementary material accompanying this paper.

2.1 Studying Character Formation with Directed Networks

In contrast to alphabetic writing systems, the Chinese characters were not created in an ad-hoc manner, but instead developed over centuries and even millennia. As a result, the formation of the Chinese characters (*zàozìfǎ* 造字法) is best described as a derivational process with some striking similarities to processes of word formation.¹ This applies specifically to the phonetic characteristics of the writing system, as reflected in the category of *xiéshēng* 諧聲 characters consisting of a phonetic element, which gives hints on the pronunciation, and a semantic element, which gives hints on the semantics encoded by a given character. Since these phonophoric characters were not formed in the same time period, their formation reflects different stages in the intertwined history of the Chinese language, which did not end with the establishment of the *kǎishū* 楷书 characters, but left its most recent traces in the character simplification carried out in the last century in China.²

As a direct result of the derivational process of character formation, Chinese characters may often reflect different stages in the history of the Chinese writing system. Similar to compound words like German *Krankheitsverlauf* ‘disease progression’, which reflects different stages of derivation and compounding, taking place at different times (see List et al. 2016a and Table 1), a Chinese character like *bǎng* 膀 ‘wing’ reflects different layers of formation, with the combination of *yuè* 月 as a semantic element with *páng* 旁 ‘side’ as a phonetic element as a late process, and the earlier confusion of *yuè* 月 and *ròu* 肉 ‘flesh’ as a character radical, and the formation of 旁 with *fāng* 方 ‘square’ as its phonetic component.³

Word	Approximate Origin	Meaning	Stage of German
kranc	900	‘weak’	Old High German
kranc	1200	‘not healthy’	Middle High German
kranc-heit	1200	‘weakness’	Middle High German
(h)loufan	700	‘run, dance’	Old High German
(h)louf	800	‘gait, course’	Old High German
Ver-lauf	1400	‘process, development’	Early New High German
Krank-heit-s-ver-lauf	1800	‘disease progression’	Modern German

Table 1: Different stages of German reflected in the compound *Krankheitsverlauf* ‘disease progression’. Data taken from the DWDS (Geyken 2010) and Pfeifer (1993).

Since Duàn Yùcái 段玉裁 (1735-1815) detected the strong correlation between the phonetic part of *xiéshēng* characters and rhymes in the Book of Odes (*Shījīng* 詩經, ca. 1050–600 BC), described in his

¹To my knowledge, this idea was first proposed by Kunze (1937) who directly compared German word formation with the formation of the Chinese characters (see also List 2008: 41-44).

²Compare, for example, *shàng* 上 as phonetic in *ràng* 让, the simplified version of 讓.

³The detailed formation of 旁 is rather complex, since many different variants for the character exist, and it would not only take too long to discuss all of them here, I am also afraid that I would lack the expertise to report the history of the character sufficiently.

Character	Middle Chinese				Phonetic (GSR)	Phonetic (SW)
方	p	j	a	ng	-	-
放	p	j	a	ng H	方	方
昉	p	j	a	ng X	方	Ø
舫	p	j	a	ng H	方	方
舫	p	j	a	ng H	方	Ø
旁	b		a	ng	方	-
謗	p		a	ng H	方	旁
謗	p		a	ng H	方	旁
膀	p		a	ng H	Ø	旁

Table 2: Contrasting Karlgren's and Xǔ Shèn's analysis of Chinese characters containing 方. Ø indicates that the character is missing in the rouse, the dash indicates that the character is not analysed phonetically.

Liùshū Yīnjùn Biǎo 六書音均表, Historical Chinese Phonology is equipped with a powerful tool for the reconstruction of Old Chinese phonology.⁴ While the rhymes only provide information on the finals of Old Chinese, *xiéshēng* characters can help to shed light on the initials as well. Since then, the reconstruction of Old Chinese phonology has greatly improved, culminating in a series of very similar reconstruction systems, which have been proposed independently by different scholars in the beginning of the 1990s (Behr 1999) and since then been constantly improved (Baxter and Sagart 2014, Sagart 1999, Schuessler 2007).

The standard way of employing *xiéshēng* characters for the purpose of Old Chinese reconstruction is to assemble characters into so-called *xiéshēng series* by which characters are clustered into groups that share the same phonetic. From the divergent character readings that can be found in Middle Chinese as reflected originally in the rhyme book *Qièyùn* 切韻 (601 AD), scholars carry out an internal reconstruction of Old Chinese character readings which explain differences in Middle Chinese based on general ideas about sound change tendencies (for details, see Baxter 1992). *Xiéshēng series* are thereby defined rather loosely. What is considered as bearing the phonetic value in a given character does not necessarily coincide with the classical practice of *liùshū* 六書 analysis as presented in character dictionaries like Xǔ Shèn's *Shuōwén Jiězì* 說文解字 (121 AD). Karlgren's *Grammata serica recensa*, which is the most prominent collection of *xiéshēng series*, for example, assigns the character *bǎng* 膀 'wing' to the *xiéshēng series fāng* 方 'square' (Karlgren 1957: #0740), while Xǔ Shèn gives *páng* 旁 'side' as the phonetic (*cóng ròu*, *páng shēng* 从肉。旁聲).

The current practice of grouping characters into rather loosely connected *xiéshēng series* in Old Chinese reconstruction often does not properly reflect the derivational aspect of the Chinese writing system. As an example, consider Table 2 where different characters which can be assigned to Karlgren's series 0740 are given in Middle Chinese reading (following Baxter 1992) along with the phonetic analysis given in the *Shuōwén Jiězì*. While Karlgren identifies the same phonetic 方 for all characters, Xǔ Shèn provides two phonetics, 方 and 旁. While one could argue that we are dealing only with different analyses here, and that Karlgren is probably right, given the similarity of the Middle Chinese readings, we can also see that those characters where Xǔ Shèn identifies 方 as the direct phonetic element belong to the third division (*sānděng* 三等) in the rhyme table system of Middle Chinese (as reflected by the medial -j-),⁵ while the characters that are given 旁 as phonetic all belong to the first division (*yīděng* 一等). This shows that the mysterious Old Chinese A/B distinction, which divides Old Chinese syllables in two distinct groups, one surfacing as

palatal rhyme in Middle Chinese, the other as non-palatal (Sagart 1999), is – at least to some degree – also reflected in the *xiéshēng* characters.

That *xiéshēng* series can be further subdivided is by no means a new discovery. It is reflected in later revisions of Karlgren’s influential system (Schuessler 2009, Zhèngzhāng 2003), but also in the recent literature, where scholars have illustrated that subseries of a given *xiéshēng* series may reflect not only the A/B distinction (Sagart and Mǎ 2017), but also further distinctions, like uvulars versus velar stops, or vowel quality (see Hill 2015: 53 and Schuessler 2009: 246f). What we lack so far, however, is a consistent way to model and analyze which *xiéshēng* series can be subdivided and how this subdivision is substantiated with respect to hypotheses on character formation.

Recalling the analogy between character and word formation mentioned above, we can try to find inspiration in the way in which word formation is handled in linguistics. A crucial aspect of word formation (but also of character formation) is the hierarchical process by which words are derived from each other at different times. If we have a compound word consisting of three base morphemes, like German *Ellenbogensgesellschaft* ‘dog-eat-dog society’ (lit. ‘elbow society we can recursively split the word into its respective components which have usually be coined during different time periods.’⁶ This analysis, which is illustrated in Figure 1, reflects a directed network of word formation processes in which the different subdivisions of the compound correspond to the nodes in the network, and the directed edges reflect the processes by indicating which part was used to compose which more complex part in the history of the formation of the compound word. The usage of networks to model word formation is quite common in linguistic morphology (Hippisley 1998, Ševčíková and Žabokrtský 2014). Given the similarity between word and character formation, it seems straightforward to also use networks to model processes of character formation in the history of the Chinese writing system.

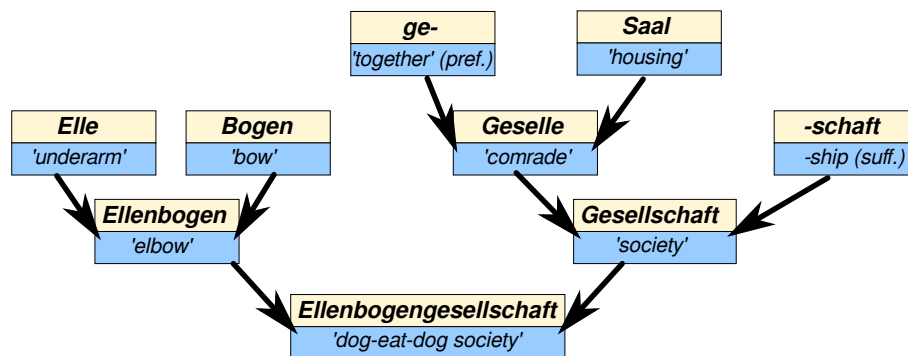


Figure 1: Word formation network for German *Ellenbogensgesellschaft*, based on information drawn from the DWDS (Geyken 2010) and Pfeifer (1993).

For initial attempts to model Chinese character formation in networks, it is useful to concentrate only on one aspect of *xiéshēng* characters, either their phonetic or their semantic component. Handling both phonetics and semantics at the same time might seem appealing at the first sight, but given that for the purpose of Old Chinese reconstruction, the semantic aspect is usually given less prominence, it should already be interesting enough if we could simply create a model based on the phonetic elements which reflects derivational history of a given *xiéshēng* series.

To illustrate how this can be done in concrete, I prepared a little experiment in which Chinese characters sharing phonetic elements are represented in a directed network. The principle by which the network is

⁴While the usefulness of using *xiéshēng* series for historical reconstruction was already detected much earlier (Behr 2005: 33-35), it was, to my best knowledge, Duàn Yùcái who first noted that *xiéshēng* characters with the same phonetic element could freely rhyme in the Shījīng.

⁵For an overview on the role of divisions, see Shen (2017) and Branner (2000).

⁶As can be seen from English *elbow*, the compound may even be traced back to the common ancestor of German and English.

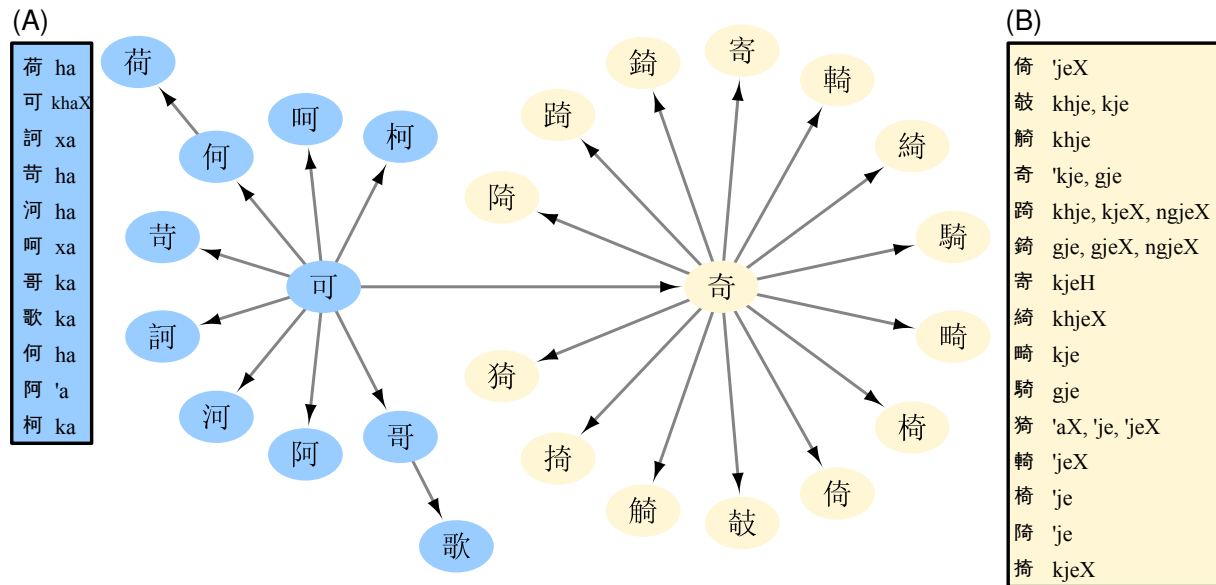


Figure 2: Character formation network of Karlgren's series #0001. Comparing the division into two clusters with the Middle Chinese character readings, one can clearly see that the A/B distinction is also reflected in this series.

constructed are straightforward. Nodes in this network are represented by characters, and phonetic relations between characters are represented by directed edges. Whenever a character *A* can be directly divided into a semantic and a phonetic component and the phonetic component constitutes a character *B* of itself, an edge is drawn from the *B* to *A*, reflecting that *A* has *B* as its phonetic determiner. For small series such a network can be easily drawn manually, but for larger amounts of data, it is useful to do this automatically, especially also, because thanks to recent digitization efforts, a large amount of data on character structures and *xiéshēng* analyses by different authors is now freely available.

Since this is a first experiment on character formation networks, I restricted the analysis to the characters and *xiéshēng* series of Karlgren (1957) which have recently been made available online in form of a WikiBook (Wikipedia Users 2017). Since Karlgren's analysis only contains information on *xiéshēng* series in general, but not on their further structure, I employed the dataset on Chinese character structures provided by Kawabata (2014). This dataset provides information on the semantic and the phonetic structure for as many as 18348 Chinese characters. Information on the phonetic component of a given character follows Xǔ Shèn's practice, by which, for example, the character *fáng* 房 is glossed as *fāngshēng* 方聲, indicating that 方 is the phonetic of this character.⁷ For 3783 out of 5142 different characters in the digital version of the GSR, a phonetic analysis could be found in Kawabata's data, corresponding to 625 *xiéshēng* series in Karlgren's analysis. From these characters, the *xiéshēng* network was constructed by adding directed edges for all characters inside a given *xiéshēng* series.

An example of the network is given in Figure 2, where series *kě* 可 #0001 is plotted along with additional character readings in Middle Chinese (about half of the full network is furthermore provided in Appendix A). What we can easily see here is that this network reflects the A/B distinction as well, with characters directly derived from 可 corresponding to the first division in Middle Chinese (group A in terms of Sagart 1999), and characters derived from *qí* 奇 reflecting the third division (group B in Old Chinese).

⁷Unfortunately, no further information on the criteria by which phonetic components were identified for each character could be found. Manually checking parts of the data, however, showed that the judgments follow largely the ones which one can also find in Xǔ Shèn.

While this example itself may already be interesting in the context of *xiéshēng* series and Old Chinese reconstruction, the real advantage of having used an automatic approach for the construction of *xiéshēng* networks is that we can now search the network automatically for similar series in which A and B are contrasted in a similar fashion. This can be done in a straightforward way by clustering each *xiéshēng* series in our graph into groups of nodes which share the same parent node (i.e. the same phonetic component). Given that the Middle Chinese readings for all characters are available, we can then check whether a given subseries shows a frequent number of third division characters or not. If one of Karlgren's series which can be further subdivided shows both subdivisions of the third division and another division, we have found a series that is similar in structure to #0001, i.e., in reflecting the A/B distinction.

Performing this experiment on the entire data, using 70% of third division readings in Middle Chinese as the threshold by which a subseries is classified as reflecting the B-group of Old Chinese, and no more than 30% of third division readings being taken as threshold to classify a subseries is classified as A-group, revealed that as many as 16 of the 625 *xiéshēng* series in the data are very likely to reflect an A/B distinction. Among these are #0001 with the subseries 可 (A) vs. 奇 (B) as discussed above, #0094 女 / 如 (A) vs. 奴 (B), which was already discussed by (Sagart and Mǎ 2017), and the above-mentioned #0740 方 (B) vs. 旁 (A). Table 3 shows eight of the 16 series which qualify as particularly striking examples. A full table showing all 16 groups is given in Appendix B, and the data as well as the source code that were used to perform this analysis are provided as part of the supplementary material accompanying this paper.

GSR	Subseries	A/B	Size	Examples
0001	可	-	11	柯何河苛呵訶阿哥奇
	奇	j	21	騎錡寄掎畸綺踦敺綺倚椅猗畸
0002	我	-	9	俄娥峨莪譌餓鷺蛾
	義	j	4	儀議蟻義
0004	也	j	6	匱池馳弛地施
	它	-	2	陀舵
0094	女	j	6	妝如汝奴要
	如	j	10	紒洳茹恕絮拏
	奴	-	14	拏帑弩怒拏拏笱驚呶悒
0256	袁	j	5	園猿轅遠
	農	-	7	銀寰懷園猿
0468	允	j	4	鈦吮沅爰
	爰	-	15	竣逡俊竣竣餽駿浚浚酸浚腹
0720	易	j	20	陽暘惕揚颺楊瘍惕惕惕惕惕惕
	湯	-	6	盪盪蕩蕩
0740	方	j	20	舫放昉枋瓶邠妨紡芳訪髣仿坊防魴衍
	旁	-	5	滂傍旁

Table 3: *Xiéshēng* series which reflect the A/B distinction (full list given in Appendix B).

It is obvious that my illustration has barely touched the potential of these analyses for Old Chinese reconstructions. Given that scholars have proposed that certain *xiéshēng* series should be subdivided for other reasons, one could use the approach presented here to search for additional distinctions which are not explicitly noted as such when relying on a loose grouping of *xiéshēng* characters into series. The analyses could also be used to enhance network approaches to rhyme analysis, following up on the work presented in List et al. (2017a). Due to the sparsity of the data in the Book of Odes, characters may easily be assigned to different rhyme groups in automatic analyses, due to the lack of evidence. By adding the phonetic character structure as an additional component, the network analysis could be strengthened and the analysis could be brought closer to classical scholarship. Even more important, however, seems the potential impact on scholarly discussion. If experts working on *xiéshēng* analyses started to provide their data in network form,

listing direct phonetic components which reflect the derivational character of the Chinese writing system, it would be much easier to compare different analyses and to build on the research of our colleagues.

2.2 Directed Networks of *fǎnqiè*

Historical glossing practices by which the pronunciation of a character was elucidated are a particularly interesting aspect of Historical Chinese Phonology. As it is well known, the Chinese writing system gives only minimal hints regarding the pronunciation of its characters. A character like 手 “hand”, pronounced as *shǒu* (or [ʃɔu²¹⁴] in the IPA), does not tell us anything about its pronunciation today and even less about its pronunciation in the past. Since the ancient Chinese scholars didn’t have an alphabet to simply transcribe their sounds (intensive contact with Indian phoneticians started much later), they started from simple equations according to which one character was pronounced similar to another character. In his *Shuōwén Jiězì*, Xǔ Shèn occasionally uses the formula “read [this character] as X” (*dúruò* X 读若《丙》) in addition to his explanations of the meanings and the structure of the characters. The disadvantage of the *dúruò* method, as linguists often call it (Coblin 1983), is that it only allows to gloss characters for which a simple character with an identical pronunciation exists. It is also not clear whether the formula consistently pointed to strictly identical pronunciations or whether certain deviations were allowed. The so-called *fǎnqiè* spellings (Branner 2000, Coblin 1983) provided a much more precise way of glossing character pronunciations. This spelling method, which seems to go back to at least the third century, is based on breaking the character pronunciation into two parts, the *initial* and the *final*, and selecting two different characters glossing, one with an identical initial sound and one with the identical final. Figure 3 contrasts the *dúruò* practice with the refined *fǎnqiè* spelling.



Figure 3: Contrasting the *dúruò* glossing and the *fǎnqiè* spelling.

Given their straightforwardness and simplicity, *fǎnqiè* pronunciation glosses became quite popular among Chinese scholars, and even today, people may occasionally use them in order to explain pronunciations without having to rely on foreign writing systems like the Latin alphabet. As a result, there is an abundance of sources which use this pronunciation device throughout the history of the Chinese language. Although the pronunciation is only given indirectly, with respect to the pronunciation traditions which were active at a given epoch, the *fǎnqiè* spellings offer great help to explore how the pronunciation of the Chinese language changed over time.

Most of this research on pronunciation changes investigated through the comparison of *fǎnqiè* spellings has been carried out manually. Chinese scholars’ systematic work on *fǎnqiè* spellings goes back to the early 19th century, when scholars like Chén Lǐ 陳澧 (1818-1882) began to investigate in his *Qìyùn kǎo* 切韻考, which characters were used to denote certain initial sounds (*fǎnqiè shàngzì* 反切上字), and which characters were used to denote the finals (*fǎnqiè xiàzì* 反切下字). As we can expect, instead of using the same character for the pronunciation of the initial sound all the time, scholars would often alternate the characters, but the alternations were more or less consistent, with some characters being used more

frequently and some characters being used less frequently. Chén Lǐ figured out that the characters could be classified in a rather rigorous manner which would allow to reconstruct direct pronunciations of the *fǎnqiè* spellings. For example, we can say that the characters *gōng* 公, *gǔ* 古, *gàn* 干, etc. were regularly used in the *Qièyùn* to indicate initials which would be spelled as [k] in the IPA, while *kǒu* 口, *kě* 可, and *kǔ* 苦 were used to pronounce [k^h].

Chén Lǐ summarized the 452 different characters that were used as initials in the *fǎnqiè* spellings of the *Qièyùn* into 40 different initial groups (*shēnglèi* 聲類). Later in 1915, Bernhard Karlgren (1889-1978) postulated a system with as many as 47 initials, 32 basic initials, and “dans 15 cas une série pure et une série yodisée” (Karlgren 1915–1926: 93), referring to palatalized consonants ([p^j, k^j, t^j]). Chao Yuenren (1892-1982), however, showed in 1941 that a class of initials does not necessarily correspond directly to an initial in Middle Chinese, since many of the distinctions turned out to be in complementary distribution with the finals, i. e., “there is a tendency, in varying degree for various initials, for the upper ch’ieh word [*fǎnqiè shàngzì*, JML] to agree with the lower.” (Chao 1941 [2006]: 309).

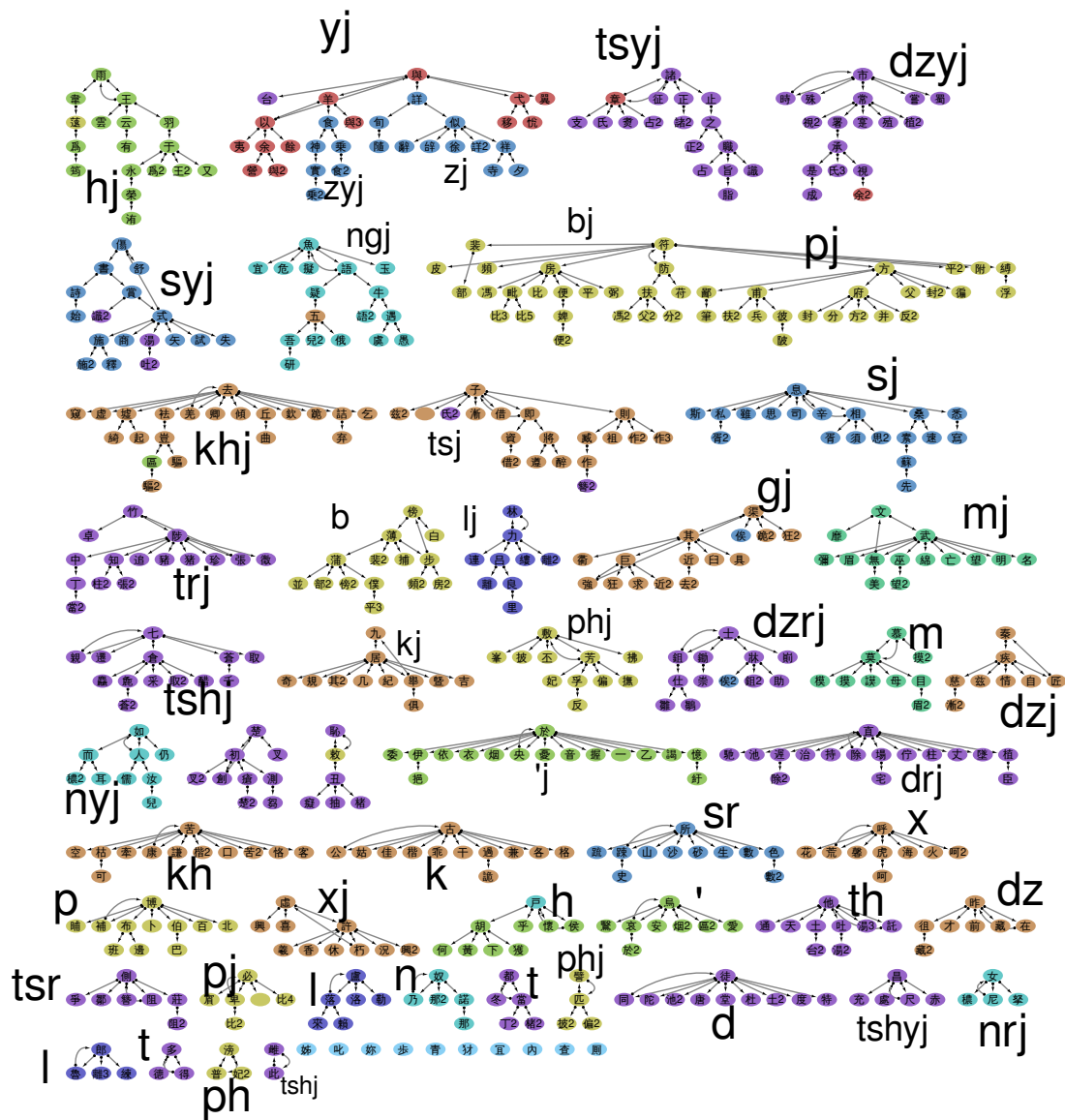


Figure 4: Full network of *fǎnqiè* initial characters in the *Guǎngyùn*. Labels indicate approximate initials along with palatal medials for a group of nodes in the network.

What is interesting about the research initiated by Chén Lǐ and its successors inherently employed rudimentary network thinking to arrive at their clusters (Gěng 2004). The system underlying the *fǎnqiè* spelling, consisting of a *glossed character* and a *glossing character* can be easily translated into a system of *directed networks* in which we draw a link from the glossing character to the glossed character. In order to illustrate how this principle can be used in concrete, I constructed a network of the *fǎnqiè* initials of the *Guǎngyùn* (ca. 1001 AD), a later edition of the *Qièyùn*, where we find *fǎnqiè* spellings for more than 20,000 characters. In this network, I only concentrated on the initial spellers, that is, the initial consonants of the language encoded in the source, and constructed a network of all internal relations among the glossing characters (characters which were not used as glossing characters themselves were deliberately ignored). Furthermore, in order to avoid that characters with multiple readings increase the amount of noise in the data, I separated characters which were glossed more than one time, labeling them A, A2, A3, etc. This resulted in a network of 568 nodes and 565 directed edges drawn among them.

The network is shown in full in Figure 4. The source code which I used to compute the network along with the data is available from the supplementary material accompanying this paper. From eyeballing, we can see that the system looks rather structured. The network is not connected but rather splits neatly into connected components. These components themselves often exhibit a hierarchical structure which indicates that there are clear preferences for the usage of glossing characters for certain sounds. As can further be seen from my annotations to the figure, the components of the network mostly neatly correspond to Middle Chinese readings provided we also look at the medial, and the network thus provides a visual account on Karlgren's *série yodisée*.

When inspecting particular groups, however, we can likewise find a couple of *transitional* groups, in which the reconstructed Middle Chinese spellings are less clearly separated. This holds specifically for the second cluster from top on the right of the figure, consisting of non-aspirated bilabial stops, where we can see that voiced and voiceless stops (*bj* and *pj*) form a transitional zone, in which it is difficult to find a clear-cut distinction between voiced and voiceless stops. Although we know that the palatalized series of bilabial stops was extremely unstable, with most of its members later merging into [f] in most dialects, this does not seem to be the only reason for the structure of the cluster in our network.

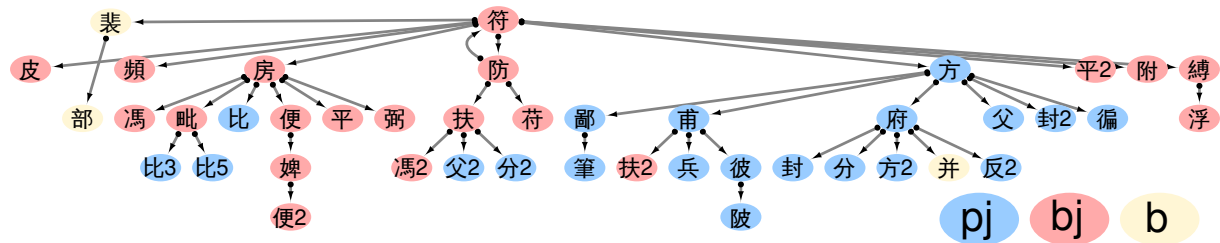


Figure 5: The bilabial “transitional” cluster from Figure 4 in higher resolution.

Unfortunately, when inspecting the subgraph of *bj* and *pj* in Figure 5, that the transition between the initials is rather an artifact of the way the network was constructed than an indicator of fuzzy *fǎnqiè* spellings due to the closeness of the sounds in Middle Chinese. What looks like a transition is a direct result of the above-mentioned decision to split characters which are glossed multiple times automatically into several nodes in the network. This procedure has the clear disadvantage that – by splitting one character into potentially multiple readings, because it is glossed several times – we do not know which of the readings is the basic one which was intended when the authors of the *Guǎngyùn* employed this character to gloss other characters. Thus, in the *Guǎngyùn*, *fú* 符 ‘a tally’ with the Middle Chinese reading *bju* clearly glosses 方, which has the normal Middle Chinese reading *pjang*. But 方 also occurs in the homophone group *bjang* in the *Guǎngyùn*. By replacing our node 方 with the node 方2, which itself is the network would be separated into two parts, and we can manually re-assign the correct Middle Chinese readings for the characters in question. I have tried to illustrate this manual correction of the automatic network in Figure 6 where those

nodes whose supposed reading has been manually adjusted are marked. It seems that the distinction between the initials *pj* and *bj* in Middle Chinese is after all consistently reflected in the *Guǎngyùn*.

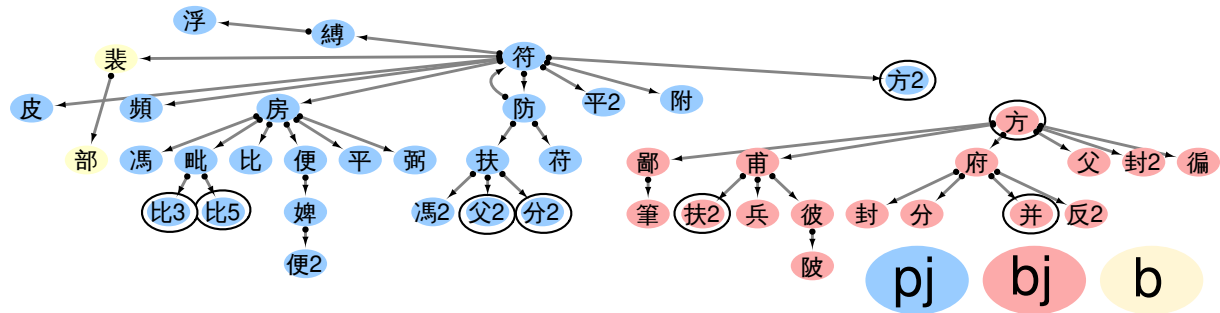


Figure 6: Manually refined network after inspection and correction of the network in Figure 5. Nodes whose assignment to a Middle Chinese reading has been modified manually are marked by a circle with black border around them.

This pilot study on network approaches to the study of *fǎnqiè* spellings may seem disappointing. So far, I could not provide a consistent solution to handle characters which are glossed differently in sources like the *Guǎngyùn*. Instead, the relevant problems resulting from the naive approach which I used still need manual refinement. I nevertheless think that it is a valuable enterprise to further explore the potential of network approaches to pronunciation glosses in Historical Chinese Phonology. Even if we have to take the time to manually refine the networks we can infer from *fǎnqiè* collections, we gain a considerable amount of transparency, since – in contrast to the classical approaches by which the results of a purely qualitative analysis are presented in tabular form – the network display is not only visually appealing but also much more transparent, in so far as it shows all variation that can be found in the data.

Networks of pronunciation glosses in Historical Chinese Phonology are still underexplored, both with respect to traditional scholarship and with respect to the way they are best handled and analyzed in quantitative approaches. They are by no means limited to *fǎnqiè* spellings in rhyme books but could be easily adapted to account for *fǎnqiè* spellings in critical editions of the classics, like the *Jīngdiǎn Shìwén* 經典釋文, as well as various forms of more direct glosses, like the *dúruò* glosses in the *Shuōwén* or the sound glosses in the *Shímíng* (Bodman 1954). The crucial question remains, how the network modeling can be further improved in order minimize the amount of qualitative analysis and manual refinement when constructing these networks.

If we could develop an approach that would infer the clusters of different sound glosses that point consistently to the same sound, they could give us fascinating insights not only into the phonological system of Chinese varieties spoken during a given time period, but perhaps also into the dynamics underlying pronunciation changes, when comparing different networks across different times and places. In order to achieve this, however, the modeling and the analysis of data on pronunciation glosses in directed networks needs to be enhanced. I hope, however, that this small illustration is sufficient to show that it is worthwhile to think along these directions when dealing with Middle Chinese and even earlier stages of Chinese.

2.3 Using Network Approaches for Evaluation

In the previous sections, I have tried to show how network approaches can help to analyze Middle Chinese and Old Chinese data for the purpose of enhancing our understanding about linguistic reconstruction in Historical Chinese Phonology. Additionally, network approaches are also useful for the *evaluation* of the analyses which have already been produced by different scholars. In the following, I will concentrate on two aspects: the comparison of reconstruction systems with help of *data-display networks*, and the comparison of rhyme analyses provided by different scholars. As all approaches presented in this paper, these analyses represent work in progress rather than final results.

2.3.1 Comparing Reconstruction Systems with Data-Display Networks

When dealing with different reconstruction systems for Old Chinese phonology, it is quite difficult, even for experienced scholars, to spot the actual differences between the systems. That these differences exist, and that they can be quite substantial, is beyond question and easy to understand if one takes into account that Old Chinese is reconstructed with help of a *philological* (as opposed to a mainly *comparative approach*) by which data from different sources is sifted and individually weighted (cf. Jarceva 1990: 409 and List 2008).

When comparing different reconstruction systems, it is not enough to simply look at the inventories of proto-phonemes proposed by different scholars. Even if two proto-inventories are exactly the same, it is possible that scholars will provide different reconstructions for individual characters. The only way to compare two or more reconstruction systems is therefore to compare the concrete reconstructions for a certain number of characters. In addition to the sample of words, however, we also need a clear account on which segments (which proto-sounds) should be compared with each other. When comparing proto-forms for Chinese *yī* ‘one’ in different Old Chinese reconstruction systems, such as Karlgren (1950) *ʔiět, Li (1971) *ʔjit, Wáng (1980) *iet, and Baxter and Sagart (2014) *ʔi[t], we would obviously not compare the medial *i in Karlgren with the initial *ʔ in Baxter and Sagart. When adding more reconstructions, such as the one for *qī* 七 ‘seven’ across the four systems, for which the authors give *ts'jëť, *tshjit, *tshiet, and *[tsʰ]i[t], respectively, we can further see that there are not only differences for the different segments in the same positions, but also for the interpretation of the words. Although all authors give different medials, main vowels, and finals in the two words, they are structurally consistent in giving both words the *same* sound segments for medial, nucleus, and coda.

What we can see from this example is that difference in the sound segments, like the choice of initials, or the concrete solution proposed for a problem like the A/B distinction between Old Chinese syllables (Sagart 1999), does not immediately reflect important differences in the reconstruction systems. If two scholars just choose another symbol for a distinction that they both recognize and acknowledge, this does not render the reconstructions *incompatible* and should therefore not be used as a criterion for dismissing a given reconstruction system, at least not in a first step. If two systems are *structurally equivalent*, they have equivalent predictive power for the descendant language(s) they are supposed to reconstruct.⁸

Although this *abstractionist* notion of proto-forms, which can already be found in the work of Saussure (1916) and Meillet (1903),⁹ is problematic for the endeavour of linguistic reconstruction and usually not strictly followed (Lass 2017), the potentially abstract notion of proto-forms is important to be kept in mind when comparing different reconstruction systems. When distinguishing the *structural differences*, resulting from the direct interpretation of the data and the identification of regular sound correspondences, from the *substantial differences*, resulting from a phonetic and phonological interpretation of the identified correspondences, we have a much clearer account on the core of the differences, and whether they are worth our consideration or not.

Wáng (1980)	tsh	i	e	t	+	-	i	e	t
Baxter and Sagart (2014)	[tsʰ]	-	i	[t]	+	?	-	i	[t]
Li (1971)	tsh	j	i	t	+	?	j	i	t
Karlgren (1954)	ts'	i	ẽ	t	+	?	i	ẽ	t

Figure 7: Exemplary alignment of different Proto-Forms with help of the EDICTOR tool.

⁸This presupposes, of course, that we model sound change as a mechanical process that does not judge the plausibility of sound shifts based on their substance (e.g., arguing that it is more plausible that a [p] becomes an [f] than *vice versa*).

⁹Compare Saussure's statement that one could likewise replace the proto-forms by numbers ('[...] qu'on pourrait désigner les éléments phoniques d'un idiome à reconstituer par des chiffres ou des signes quelconques.' Saussure 1916:303) and Meillet's opinion, according to which the proto-forms only serve as a convenient way to refer to the sound correspondences ('Les «restitutions» ne sont rien autre chose que les signes par lesquels on exprime en abrégé les correspondances' Meillet 1903: 42).

But how can we compare reconstruction systems *structurally*? Firstly, we need to have the data assembled in *aligned form*, in order to make sure that we only compare like with like (e.g., medial with medial, and vowel with vowel). A sample illustration in which alignments of the proto-forms for ‘seven’ and ‘one’ were produced with help of the EDICTOR tool (List 2017a) is given in Figure 7. Alternatively, we can also select a single aspect, such as, for example, the vowel system proposed in different reconstruction systems. Having assembled a substantial amount of different proto-forms in this way, the structural comparison between different reconstruction systems can be modeled as a comparison of different *cluster analyses*, or, more accurately, *partitioning analyses*. A partitioning analysis assigns a given number of objects to a certain number of different groups. When dealing only with the vowels proposed by different reconstruction systems, we can say that a given reconstruction, like the one by Karlgren, for example, assigns each Chinese character, for which a proto-form is given, to a particular group depending on the main vowel selected for the reconstruction.

If, for a given number of reconstructions, we model each reconstruction system as a partitioning analysis, based on the main vowel proposed by the system, we can use standard metrics from graph theory and Natural Language Processing to compare different reconstruction systems with each other. A very straightforward measure for the comparison of two partitioning analyses are the so-called B-Cubed scores (Amigó et al. 2009), which have proven specifically useful for the evaluation of automatic cognate detection methods in historical linguistics compared to a gold standard (Hauer and Kondrak 2011, List et al. 2017a). Being an evaluation measure, B-Cubed scores come in the typical three flavors of *precision*, *recall*, and *F-Score*, with precision being similar to the notion of *true positives*, and recall to *true negatives*. For the purpose of comparing reconstruction systems, only the *F-score* is needed, as it is a symmetric measure, and the notion of true positives and true negatives is senseless, unless we decide that we blindly trust one of the given systems. As the scores for precision and recall, the F-score score ranges between 0 and 1, with 1 indicating that two partitioning analyses are identical.

In order to compare more than one reconstruction system, we can make use of techniques for *exploratory data analysis* (Morrison 2014), such as the NeighborNet algorithm (Bryant and Moulton 2004) as provided by the SplitsTree package (Huson 1998), which tests the tree-likeness of a given classification and is also very popular in linguistics, where it is used to test for the degree of reticulation among different language varieties in dialectology and historical language comparison (Bryant et al. 2005, Heggarty et al. 2010). People often refer to the networks produced by the NeighborNet algorithm as *phylogenetic networks*. This, however, is a misnomer, as the algorithm does not provide a true evolutionary analysis (and it was never supposed to do so), but instead provides a view on the data, displaying obvious similarities and differences among the different objects of investigation. They are therefore better called *data-display networks* (Morrison 2011) and should be distinguished from *evolutionary networks* which are supposed to display how a given set of entities evolved in a network-like fashion, allowing for different kinds of reticulation.

In order to illustrate how data-display networks can be used to study differences among Old Chinese reconstruction systems, I designed a little experiment, based on data taken from (List et al. 2017b), who provide Old Chinese reconstructions for all rhyme words in the Shījīng for eight different reconstruction systems (Baxter and Sagart 2014, Karlgren 1950, Li 1971, Pān 2000, Schuessler 2007, Starostin 1989, Wáng 1980, Zhèngzhāng 2003). In order to keep the analysis simple, I only extracted the different reconstructions of the main vowel for each character in each system, and carried out a pairwise comparison of all eight systems, computing the B-Cubed F-scores for each pair, omitting characters for which no reconstruction could be found in the data. These scores were then converted to a distance matrix and fed to the NeighborNet algorithm (the source code is provided along with this paper).

The resulting network is provided in Figure 8. As one can see, the data roughly clusters in three larger subgroups, namely Schuessler, Baxter and Sagart, and Starostin vs. Pān and Zhèngzhāng vs. Karlgren, Li, and Wáng. On a larger scale, we can divide the data into all six-vowel systems versus the non-six-vowel systems (Karlgren, Wáng, Li). Given that Pān is a direct student of Zhèngzhāng, the closeness between their reconstruction systems is not directly surprising. What may be surprising is the closeness of Schuessler,

Starostin, and Baxter and Sagart, given the notable differences between the systems with respect to the criterion of *vowel purity* tested in List et al. (2017b). Even if the network analysis cannot directly explain all these differences in detail, it seems like a worthwhile enterprise which should be further expanded by comparing not only the vowels, but fully aligned proto-forms. Given the straightforwardness of the application, it seems also useful to test it on other language families, where there is similar disagreement as in the reconstruction of Old Chinese phonology.

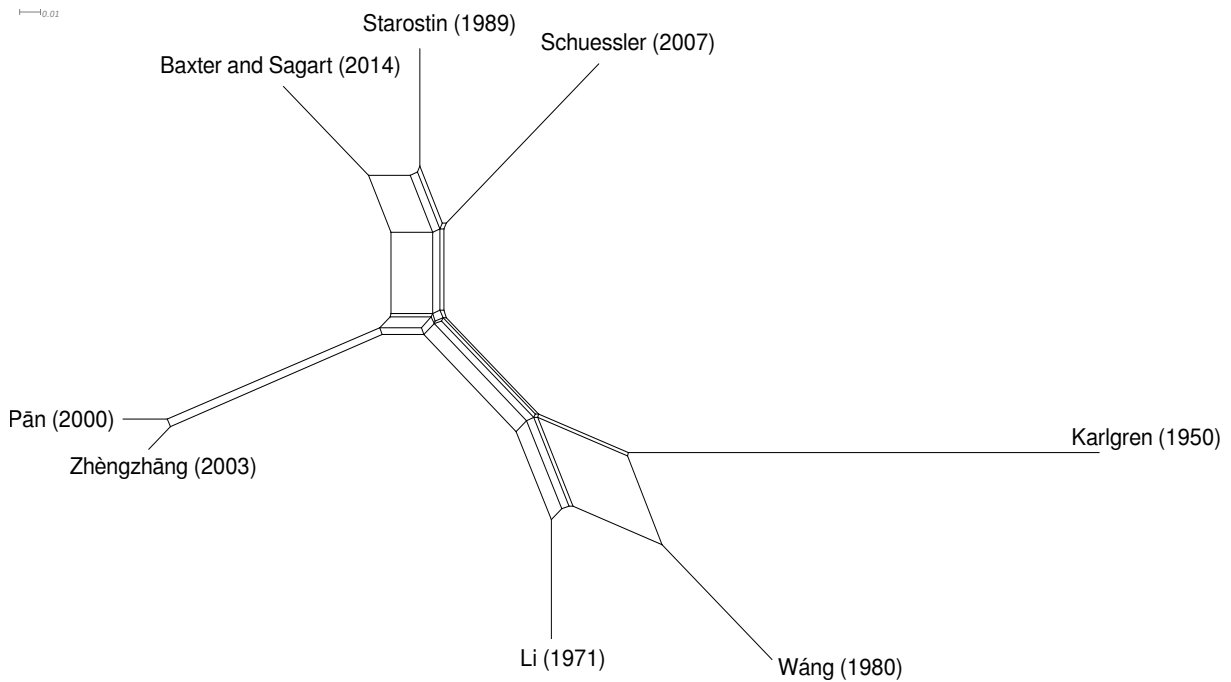


Figure 8: NeighborNet of eight different reconstruction systems for Old Chinese characters from the Shījīng

2.3.2 Comparing Differences in Rhyme Annotations

In List et al. (ibid.), rhyme networks were used to test to which degree different reconstruction systems conform to what Ho (2016) calls ‘vowel purity’, namely the hypothesis that rhyming practice in Old Chinese (and probably also later) was very strict in adhering to identical vowels in rhyming. The test by List et al. (2017b) revealed that the system of Baxter and Sagart (2014) (and of six-vowel theories of Old Chinese in general) is reflecting the principle of vowel purity much more closely than do systems with more vowels (Karlgren 1950) or less vowels (Li 1971, Wáng 1980).

In this context, it is important to recall that – what was also mentioned in the paper by List et al. (2017b), but might easily be misunderstood by readers – the adherence to vowel purity cannot be used to prove or disprove a given reconstruction system, since the adherence to vowel purity is a hypothesis about Old Chinese rhyming practice itself, and we know well that vowel purity in rhyming can be easily abandoned or disregarded across rhyming traditions in different cultures. Apart from the problem that studies on vowel purity do not bear any diagnostic value with respect to the accuracy of reconstruction systems, one additional problem in the study by List et al. (2017a) is the fact that vowel purity itself was only tested by comparing the rhyme judgments of one source (Baxter 1992) with different reconstruction systems. Given that Baxter himself is reconstructing a six-vowel system on the basis of rhyme evidence, it is quite likely that the rhyme decisions proposed by Baxter could have influenced the analysis.

While alternative rhyme judgments were not available when drafting the original study on vowel purity, we have now had the time to digitize the rhyme judgments reported in Wáng (1980).¹⁰ Given that two different rhyme analyses have been digitized now, it is interesting and also important for the reconstruction of Old Chinese Phonology to check to which degree different scholars differ in what they judge to rhyme and what not.

We can think of different measures to compare the difference in the actual rhyme judgments of the two versions. A simple measure is to compare, how many stanzas differ. From 1014 common stanzas,¹¹ 131 are different between Wáng and Baxter, which amounts to 12.9%. A far more interesting aspect is to check *how much* different stanzas differ. Similar to the partitioning task in comparing actual linguistic reconstruction systems for Old Chinese, exemplified in Section 2.3 above, we can do this with help of the B-Cubed scores, since the assessment for a given stanza, whether to words rhyme or not, can also be thought of as a clustering task (authors decide which words belong to the same rhyme partition in a given cluster). Applying B-Cubed scores to compare the rhyme judgments, we find 96.8% of similarity between Baxter's and Wáng's rhyme judgments. This means that the internal difference between the rhyme judgments by Baxter and Wáng is less pronounced than one might think when only checking whether a given stanza is interpreted differently in *any* way.

Following up on the Shījīng Rhyme Browser application by List (2017b), a new rhyme browser has been created, which is this time specifically dedicated to illustrating the differences in the rhyme judgments by Baxter and Wáng. It can be accessed at <http://digling.org/shijing/wangli/>. In the future, we hope to find time to add more rhyme analyses by different scholars to the sample. What complicates this procedure, however, is that only a few scholars have so far annotated their rhyme judgments of the Book of Odes in a fashion similar to Baxter and Wáng. If rhyme judgments are not provided in transparent form, this does not only exacerbate the task of digitizing the analyses, it also bears the danger of misinterpreting what scholars originally meant.

3 Summary and Outlook

All approaches that I have presented in this paper need to be considered as strict work-in-progress. Currently, none of the applications would be sufficient to write a paper on it by itself. Nevertheless, since I believe that these ideas may prove useful in future research, I think that it is justified, if not even needed, to present them in this context. On the one hand, I hope that they can initiate a discussion among experts of Historical Chinese Phonology that may help to improve the methods in the future. Second, I hope that my examples can convince scholars of both Historical Chinese Phonology and computational disciplines that it is worthwhile to rethink the way we compare and analyze data in Chinese historical linguistics at the moment. Quantitative methods as I have presented them here do not aim at replacing scholarly expertise by machines. Instead, by rendering the knowledge that has been accumulated during the past centuries more transparent, they may guide experts in Historical Chinese Phonology in testing existing and in creating new hypotheses. Given how straightforward it is to restate some of the classical problems in the reconstruction of Middle and Old Chinese as network problems for which quantitative solutions are available further illustrates that Historical Chinese Phonology has always been an inherently data-driven discipline. I think it is time that we embrace this aspect by working on integrated solutions in which qualitative and quantitative analysis go hand in hand.

¹⁰This is common work with Nathan W. Hill for a project that we currently label the Chinese Historical Phonology Database (CHIP), but whose name may easily change in the future.

¹¹We excluded multi-morpheme rhymes.

References

- Amigó, E., J. Gonzalo, J. Artiles, and F. Verdejo (08/2009). “A comparison of extrinsic clustering evaluation metrics based on formal constraints” . *Information Retrieval* 12.4, 461–486.
- Baxter, W. H. (1992). *A handbook of Old Chinese phonology*. Berlin: de Gruyter.
- Baxter, W. H. and L. Sagart (2014). *Old Chinese. A new reconstruction*. Oxford: Oxford University Press.
- Behr, W. (1999). *Odds on the Odes*. PDF: <https://web.archive.org/web/20110519085435/http://www.ruhr-uni-bochum.de/gpc/behrr/HTML/Excellence.htm>. (Visited on 04/27/2017).
- (2005). “Language change in premodern China: Notes on its perception and impact on the idea of a “constant way”” . In: *Historical truth, historical criticism and ideology. Chinese historiography and historical culture from a new comparative perspective*. Ed. by H. Schmidt-Glintzer, A. Mittag, and J. Rüsen. Leiden and Boston: Brill, 13–51.
- Ben Hamed, M. and F. Wang (2006). “Stuck in the forest: Trees, networks and Chinese dialects” . *Diachronica* 23, 29–60.
- Bodman, N. C. (1954). *A linguistic study of the Shih Ming. Initials and consonant clusters*. Cambridge: Harvard University Press.
- Branner, D. P. (2000). “The rime-table system of formal Chinese phonology” . In: *History of the language sciences. An international handbook on the evolution of the study of language from the beginnings to the present*. Ed. by S. Auroux, E. F. K. Koerner, H.-J. Niederehe, and K. Versteegh. Vol. 1. 18. Berlin and New York: de Gruyter, 46–55.
- Bryant, D. and V. Moulton (2004). “Neighbor-Net. An agglomerative method for the construction of phylogenetic networks” . *Molecular Biology and Evolution* 21.2, 255–265.
- Bryant, D., F. Filimon, and R. D. Gray (2005). “Untangling our past: Languages, Trees, Splits and Networks” . In: *The evolution of cultural diversity: A phylogenetic approach*. Ed. by R. Mace, C. J. Holden, and S. Shennan. London: UCL Press, 67–84.
- Chao, Y. (1941 [2006]). “Distinctions within Ancient Chinese” . In: *Linguistic essays by Yuenren Chao*. Ed. by Z.-J. Wu and X.-n. Zhao. Originally published in HJAS, 5.3/4 (1941), 203–233). Běijīng 北京: Shāngwù 商务, 304–346.
- Coblin, W. S. (1983). *A Handbook of Eastern Han Sound Glosses*. Chicago: The Chinese University Press.
- Geyken, A., ed. (2010). *Digitales Wörterbuch der deutschen Sprache DWDS. Das Wortauskunftssystem zur deutschen Sprache in Geschichte und Gegenwart*. URL: <http://dwds.de>.
- Hauer, B. and G. Kondrak (2011). “Clustering semantically equivalent words into cognate sets in multilingual lists” . In: *Proceedings of the 5th International Joint Conference on Natural Language Processing*. (Chiang Mai, Thailand, 11/08–11/13/2011). AFNLP, 865–873.
- Heggarty, P., W. Maguire, and A. McMahon (2010). “Splits or waves? Trees or webs? How divergence measures and network analysis can unravel language histories” . *Philos. Trans. R. Soc. Lond., B, Biol. Sci.* 365.1559, 3829–3843.
- Hill, N. W. (2015). “Proposal for a transcription of Chinese characters in the study of early Chinese language and literature” . *Bulletin of Chinese Linguistics* 8, 48–60.
- Hill, N. W. and J.-M. List (2017). “Challenges of annotation and analysis in computer-assisted language comparison: A case study on Burmish languages” . *Yearbook of the Poznań Linguistic Meeting* 3.1, 47–76.
- Hippisley, A. (1998). “Indexed stems and Russian word formation: A network morphology account of Russian personal nouns” . *Linguistics Faculty Publications* 36 (6), 1093–1124.
- Ho, D.-a. (2016). “Such errors could have been avoided. Review of “Old Chinese: A new reconstruction”. by William H. Baxter and Laurent Sagart” . *Journal of Chinese Linguistics* 44.1, 175–230.
- Huson, D. H. (1998). “SplitsTree: analyzing and visualizing evolutionary data” . *Bioinformatics* 14.1, 68–73.
- Jarceva, V. N., ed. (1990). *Lingvističeskij ěnciklopedičeskij slovar (Linguistical encyclopedical dictionary)*. Moscow: Sovetskaja Enciklopedija.
- Karlgren, B. (1915–1926). *Études sur la phonologie Chinoise*. 4 vols. Archives d’études orientales 15. Leiden and Stockholm: Norstedt, 1–316.
- (1950). *The Book of Odes. Chinese text, transcription and translation*. Stockholm: Museum of Far Eastern Antiquities.
- (1957). “Grammata serica recensa” . *Bulletin of the Museum of Far Eastern Antiquities* 29, 1–332.
- Kawabata, T. (2014). *Ideographic Description Sequence Data*. The project provides data on CJK ideographs in various forms. URL: <https://github.com/cjkvi/cjkvi-ids>.
- Kunze, R. (1937). *Bau und Anordnung der chinesischen Zeichen. Oder: Wie lernen wir leichter Zeichen lesen?* [Structure and assembly of Chinese characters. Or: How can we learn to read characters more easily?] Tokyo: Deutsche Gesellschaft für Natur und Völkerkunde Ostasiens.

- Lass, R. (2017). “Reality in a soft science: the metaphonology of historical reconstruction”. *Papers in Historical Phonology* 2.1, 152–163.
- Li Fang-kuei 李方桂 (1971). “Shànggǔyīn yánjiū 上古音研究 [Studies on Archaic Chinese phonology]”. *Qīnghuá Xuébào* 清華學報 9.1-2, 1–60.
- List, J.-M. (2008). “Rekonstruktion der Aussprache des Mittel- und Altchinesischen. Vergleich der Rekonstruktion-smethoden der indogermanischen und der chinesischen Sprachwissenschaft [Reconstruction of the pronunciation of Middle and Old Chinese. Comparison of reconstruction methods in Indo-European and Chinese linguistics]”. Magister thesis. Berlin: Freie Universität Berlin. PDF: <http://hprints.org/docs/00/74/25/52/PDF/list-2008-magisterarbeit.pdf>.
- (2015). “Network perspectives on Chinese dialect history”. *Bulletin of Chinese Linguistics* 8, 42–67.
 - (2017a). “A web-based interactive tool for creating, inspecting, editing, and publishing etymological datasets”. In: *Proceedings of the 15th Conference of the European Chapter of the Association for Computational Linguistics. System Demonstrations*. Valencia: Association for Computational Linguistics, 9–12.
 - (2017b). “Using network models to analyze Old Chinese rhyme data”. *Bulletin of Chinese Linguistics* 9.2, 218–241.
- List, J.-M., S. Nelson-Sathi, W. Martin, and H. Geisler (2014). “Using phylogenetic networks to model Chinese dialect history”. *Language Dynamics and Change* 4.2, 222–252.
- List, J.-M., J. S. Pathmanathan, P. Lopez, and E. Baptiste (2016a). “Unity and disunity in evolutionary sciences: process-based analogies open common research avenues for biology and linguistics”. *Biology Direct* 11.39, 1–17.
- List, J.-M., P. Lopez, and E. Baptiste (2016b). “Using sequence similarity networks to identify partial cognates in multilingual wordlists”. In: *Proceedings of the Association of Computational Linguistics 2016 (Volume 2: Short Papers)*. Association of Computational Linguistics. Berlin, 599–605.
- List, J.-M., S. J. Greenhill, and R. D. Gray (01/2017a). “The potential of automatic word comparison for historical linguistics”. *PLOS ONE* 12.1, 1–18.
- List, J.-M., J. S. Pathmanathan, N. W. Hill, E. Baptiste, and P. Lopez (2017b). “Vowel purity and rhyme evidence in Old Chinese reconstruction”. *Lingua Sinica* 3.1, 1–17.
- Meillet, A. (1903). *Introduction à l'étude comparative des langues indo-européennes*. Paris: Hachette.
- Morrison, D. A. (2011). *An introduction to phylogenetic networks*. Uppsala: RJR Productions.
- Morrison, D. A. (2014). “Phylogenetic networks: a new form of multivariate data summary for data mining and exploratory data analysis”. *WIREs Data Mining and Knowledge Discovery*.
- Newman, M. E. J. (2010). *Networks. An Introduction*. Oxford: Oxford University Press.
- Norman, J. (2003). “The Chinese dialects. Phonology”. In: *The Sino-Tibetan languages*. Ed. by G. Thurgood and R. J. LaPolla. London and New York: Routledge, 72–83.
- Pfeifer, W., comp. (1993). *Etymologisches Wörterbuch des Deutschen*. 2nd ed. 2 vols. Berlin: Akademie. URL: <http://www.dwds.de/>.
- Pān Wùyún 潘悟云 (2000). *Hànyǔ lìshǐ yīnyùnxué* 汉语历史音韵学 [Chinese historical phonology]. Shànghǎi 上海: Shànghǎi Jiàoyù 上海教育.
- Qiú Xīguī 裘錫圭 (1988 [2007]). *Wénzìxué gàiyào* 文字學概要 [Foundations of graphemics]. Běijīng: Shāngwù 商務.
- Sagart, L. (1999). *The Roots of Old Chinese*. Amsterdam: John Benjamins.
- Sagart, L. and Mǎ Kūn 馬坤 (2017). *Xiānqín shíqī xiéshēng shēngfú de xuǎnzé wèntí* 先秦時期諧聲聲符的選擇問題. University of Macau. URL: <http://www.academia.edu/35852895/>.
- Saussure, F. de (1916). *Cours de linguistique générale*. Ed. by C. Bally. Lausanne: Payot; German translation: – (1967). *Grundfragen der allgemeinen Sprachwissenschaft*. Trans. from the French by H. Lommel. 2nd ed. Berlin: Walter de Gruyter & Co.
- Schuessler, A., comp. (2007). *ABC Etymological dictionary of Old Chinese*. Honolulu: University of Hawai'i Press.
- (2009). *Minimal Old Chinese and Later Han Chinese. A companion to Grammata Serica*. Honolulu: University of Hawai'i Press.
- Shen, R. (2017). “Děng 等 (Division and Rank)”. In: *Encyclopedia of Chinese language and linguistics*. Ed. by R. Sybesma. Vol. 2. First published online in 2015. Leiden and Boston: Brill, 13–20.
- Starostin, S. A. (1989). *Rekonstrukcija drevnekitajskoj fonologičeskoj sistemy* (Reconstruction of the phonological system of Old Chinese). Moscow: Nauka.
- Wikipedia Users (2017). *Character List for Karlgren's GSR*. URL: https://en.wikibooks.org/w/index.php?title=Character_List_for_Karlgren%27s_GSR&oldid=3229525 (visited on 09/01/2017).
- Wáng Lì 王力 (1980). *Shījīng Yùndú* 詩經韻讀 [Rhyme readings in the Book of Odes]. Shànghǎi 上海: Shànghǎi Gǔjī 上海古籍.

- Zhèngzhāng Shàngfāng 郑张尚芳 (2003). *Shànggǔ yīnxì* 上古音系 [Old Chinese phonology]. Shànghǎi 上海: Shànghǎi Jiàoyù 上海教育.
- Ševčíková, M. and Z. Žabokrtský (2014). *Word-Formation Network for Czech*. Ed. by N. Calzolari, K. Choukri, T. Declerck, H. Loftsson, B. Maegaard, and J. Mariani. European Language Resources Association.
- 耿振生, G. Z. (2004). *20 shìjì Hànyǔ yǔyīnxué fāngfǎ lùn* 20 世纪汉语音韵学方法论 [20th century's methods in traditional Chinese phonology]. Běijīng 北京: Běijīng Dàxué 北京大學.

Sources

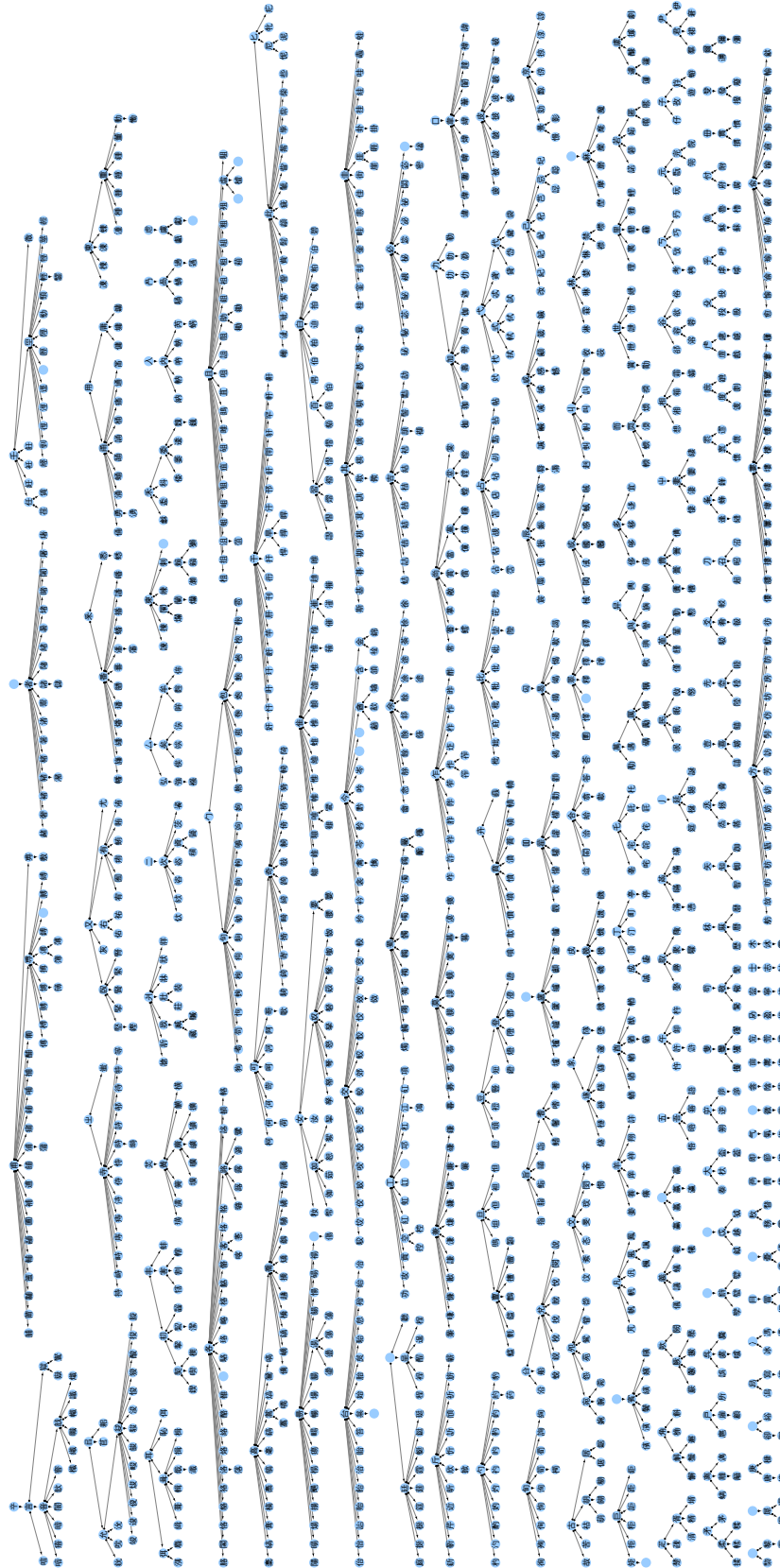
- Qièyùn kǎo* 切韵考 [Studies on the Qièyùn] (1842). By Chén Lǐ 陳澧.
- Guǎngyùn* 廣韻 [Broad rhymes] (1007/1008). By Chén Péngnién 陳彭年 and Qiū Yōng (邱雍).
- Liùshū Yīnjùn Biǎo* 六書音均表 [Phonetic tables of the six methods of character formation]. By Duàn Yùcái 段玉裁 (1735–1815).
- Shì míng* 釋名 [Explanation of names] (200). By Liú Xī 刘熙.
- Jīngdiǎn Shìwén* 經典釋文 [Explanations of the classics]. By Lù Dé míng 陸德明 (556–627).
- Qièyùn* 切韵 (601). By Lù Fǎyán 陸法言.
- Shuōwén Jiězì* 說文解字 [Explanation of simple and analysis of complex characters] (121 AD). By Xǔ Shèn 許慎 (58–157 AD). Hàndiǎn: <http://www.zdic.net>.

Supplementary Material

The supplementary material accompanying this paper contains the source code and the data as well as additional information regarding the software required to run the analyses. It is currently hosted on GitHub at <https://github.com/digling/network-in-hcp-paper>. A first version of software and source code has been submitted to Zenodo at <https://doi.org/10.5281/zenodo.1171901>.

Appendix A: Part of the Full *xiéshēng* Network

The following figure shows about half of the characters in the *xiéshēng* network presented in Section 2.1.



Appendix B: GSR Subseries with A/B Distinction

The following table is a supplement to Table 3 and lists all *xiéshēng* series which reflect the A/B distinction in their subseries structure consistently. The column A/B indicates the major character of a given subseries. If more than 70% of all characters can be assigned to the third division in Middle Chinese, this is indicated by a *j*. If less than 30% of the characters occur in the third division, this is labeled as -. In all other cases, a question mark indicates the uncertainty.

GSR	Subseries	A/B	Size	Examples
0001	可	-	11	柯何河苛呵訶阿哥奇
	奇	j	21	騎錡寄倚畸綺踦岐倚椅猗鞫
0002	我	-	9	俄娥峨莪譏餓鴛蛾
	義	j	4	儀議蟻義
0004	也	j	6	匱池馳弛地施
	它	-	2	陀舵
0049	古	?	6	故苦祜胡居
	胡	-	3	葫糊
	居	j	2	倨鋸
	固	-	3	銅個涸
0094	女	j	6	妝如汝奴耍
	如	j	10	袂洳茹絮拏
	奴	-	14	拏帑弩怒拏弩筴驚呶愀
0256	袁	j	5	園猿轅遠
	袁	-	7	鏗寰懷園猓
0468	允	j	4	鈇吮沅爰
	爰	-	15	竣逡俊燠峻餽駿浚俊酸掄
0720	易	j	20	陽暘煬揚颺楊瘍暘暢場腸湯錫惕惕湯
	湯	-	6	盪盪蕩蕩
0727	月	j	7	膾斨戕壯牀狀牂
	壯	j	3	莊裝
	將	j	5	獎漿蔣醬鏘
	臧	-	3	藏臧
0740	方	j	20	舫放昉枋旂妨紡芳訪髣仿坊房防魴防
	旁	-	5	滂傍旁
0742	亡	-	6	盲眈蝨氓育盍
	罔	j	2	網惘
0833	丁	-	6	成頂汀町亨
	正	j	2	証鉦
0835		-	5	廷呈聽
	廷	-	13	庭挺挺筵筵霆蜓鋌斑
	呈	j	6	程程醒逞程
0868	束	j	4	策策
	責	-	3	噴簣績
0918	弋	?	7	杙式忒資代
	式	j	4	拭軾試試
	代	-	4	貸岱黛袋
1041	丂	-	5	考攷巧
	号	j	2	枵鴉