

Using Mixed Incentives to Document Xi'an Guanzhong

Zhan, Juhong; Jiang, Yue; Cieri, Christopher; Liberman, Mark; Yuan, Jiahong; Chen, Yiya; Scharenborg, Odette

Publication date

2022

Document Version

Final published version

Published in

2nd Workshop on Novel Incentives in Data Collection from People

Citation (APA)

Zhan, J., Jiang, Y., Cieri, C., Liberman, M., Yuan, J., Chen, Y., & Scharenborg, O. (2022). Using Mixed Incentives to Document Xi'an Guanzhong. In J. Fiumara, C. Cieri, M. Liberman, & C. Callison-Burch (Eds.), *2nd Workshop on Novel Incentives in Data Collection from People: Models, Implementations, Challenges and Results, NIDCP 2022 - Proceedings at LREC 2022 Workshop - Language Resources and Evaluation Conference* (pp. 32-37). (2nd Workshop on Novel Incentives in Data Collection from People: Models, Implementations, Challenges and Results, NIDCP 2022 - Proceedings at LREC 2022 Workshop - Language Resources and Evaluation Conference). European Language Resources Association (ELRA).

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

Using Mixed Incentives to Document Xi'an Guanzhong

Juhong Zhan*, Yue Jiang*, Christopher Cieri†, Mark Liberman†,
Jiahong Yuan*, Yiya Chen*, Odette Scharenborg‡

*Xi'an Jiaotong University, †University of Pennsylvania, Linguistic Data Consortium,

*Baidu Research, ★Leiden University, ‡Delft University of Technology

{Corresponding Author: ccieri@ldc.upenn.edu}

Abstract

This paper describes our use of mixed incentives and the citizen science portal LanguageARC to prepare, collect and quality control a large corpus of object namings for the purpose of providing speech data to document the under-represented Guanzhong dialect of Chinese spoken in the Shaanxi province in the environs of Xi'an.

Keywords: under-resourced languages, language resources, linguistic data, annotation, novel incentives

1. Introduction

The impetus for this effort was the intersection of interests among the authors: 1) to document the Guanzhong dialect, 2) to develop low-cost, portable, scalable, and replicable language resource development methodologies, 3) to augment the supply of language resources through the use of novel incentives and 4) to innovate human, automated and hybrid methods for rapid documentation of under-resourced languages. This paper describes the Guanzhong dialect, prior efforts to create language resources to study this variety, the research goals of the current collection, the use of a citizen science platform and its implications, recruitment, incentives, collection, quality control and the resulting data.

2. Xi'an Guanzhong

A prior successful collaboration between the Xi'an Jiaotong University and the Linguistic Data Consortium, described in §3, that has resulted in published language resources (Jiang, et al. 2020), encouraged us to continue and expand the collaboration focusing on the Guanzhong dialect spoken in and around Xi'an.

2.1 Background

The Guanzhong dialect of Mandarin (hereafter Guanzhong dialect), also known as Qin language, is a dialect family spoken in the Guanzhong region of Shaanxi Province. The Guanzhong region includes five major cities: Xi'an, Baoji, Xianyang, Weinan, Tongchuan and the Yangling district. Based on the Shaanxi Statistical Report (2021), the region has a total area of 55,623 square kilometers and a population of 25,875,539. The Guanzhong dialects can be classified into two sub-dialect groups: the East-fu dialect, spoken in Xi'an and cities to its east, and West-fu dialect, spoken in e.g. Baoji and regions to the west of Xian (Li, 1989). With an estimated more than 50,000,000 speakers, it ranks among the top 40 languages, above e.g. Polish and Yoruba, in terms of the total number of speakers.

The Guanzhong dialect was once the official language of the Zhou, Qin, Han and Tang dynasties in Chinese history, and in ancient times it was called "Yayan", or the "elegant dialect".

Xi'an is the place where the Guanzhong dialect originated and developed through history. It is the central area of Guanzhong Plain and also the capital city of Shaanxi Province. Known as Chang'an, China's capital during the Tang dynasty. It was authorized by the UNESCO as a world-famous historic city in 1981. Xi'an has served as an imperial capital since ancient times and lasted through thirteen dynasties, roughly two thousand years. Chinese culture, language and writing were all formed and developed during this period.

Deeply rooted in the traditional culture of Xi'an, the Guanzhong dialect has a relatively long history and enjoys a large number of native speakers. Many ancient dialect words still remain in the Guanzhong dialect (Zhao, 2020). The traditional Shaanxi Local opera Qinqiang is sung in the Guanzhong dialect, too.

Documentation of the Guanzhong dialect dates at least to the Han dynasty in the first century BCE with Yang Xiong's development of the *Fangyan* dictionary of regional varieties. Bai (1954) produced a more recent survey of Guanzhong dialects. Current research on the Guanzhong dialect focuses mainly on phonological variations (Wang, 1995, Zhang, 2005), use of special words (Li, 2014) and personal pronouns (Sun, 2021). Despite the long tradition among Sinologists of studying the Guanzhong dialect, it has received less attention than might be expected given the number of speakers. So far, no general purpose speech corpus of Guanzhong dialect has been built, although Xing (2014) proposed the necessity of building a large dialect speech corpus.

2.2 Phonological Characteristics

Lexical tones in Guanzhong Mandarin dialects seem to have quite systematic correspondences to tones in Standard Chinese. In some varieties, the following mappings apply:

1. Yinping (the first or level tone) changes to Qingsheng (the fifth or Neutral tone) (Lu, 2010); for example *shēng chǎn* is *sheng chàn* in the Guanzhong dialect.
2. Shangsheng (the third or falling-rising tone) changes to Qusheng (the fourth or falling tone); *lǐ jiě* is *lì jiè* in the dialect.

3. Qusheng (the fourth or falling tone) changes to Yinping (the first or level tone); for example, wén jiàn is wén jiān in dialect.
4. Yangping (the second or rising tone) remains the same (Zhao, 2017). The dialectal tone of liáng pí is the same as that in Mandarin.

Liu et al. (2020) have also provided a detailed acoustic study of the tonal mapping between Xi'an Mandarin and Standard Chinese.

2.3 Lexical Characteristics

The lexical system is an important aspect of the uniqueness of the Guanzhong dialect (Li, 2014). According to current research findings, in the Guanzhong lexicon, 21.3% of the words are typically dialectal and the remaining 78.7% are consistent with or close to Mandarin (Wang, 2015). The Guanzhong dialect contains complex lexical variations in different aspects, such as pronouns, indicative pronouns, modal particle and so on (Zhao, 2020).

There are a great number of variations in the use of parts of speech. Hóu (monkey) is a noun in Standard Mandarin, but is also used as an adjective in Guanzhong dialect. For example, people could say “Zhè (this) Wā (kid) Hóu (monkey) dì (得, an auxiliary word) Tāi (very)” which means “this kid is overactive and naughty”. The word Chè (Chě in Mandarin) is a verb meaning “tear”. In Guanzhong Dialect, it can be used as an adjective in sentence like “Kū zi Chè le”, which means the trousers are torn.

3. Prior Corpus Efforts

Several of the authors had collaborated previously to begin creating corpora for the Guanzhong dialect, initially within the *Global TIMIT* framework. Global TIMIT aims to create corpora in multiple languages that share key features with the original TIMIT Acoustic-Phonetic Continuous Speech Corpus (Garofalo et al., 1993), designed to also support the development of speech to text systems. These key features included:

- many speakers
- many fluently-read sentences containing a representative sample of linguistic patterns
- time-aligned lexical and phonetic transcripts
- three classes of sentences according to whether they were read by all speakers, a few speakers, or just one speaker

Global TIMIT differs from the original in features that have proven expensive to implement or not strictly necessary or both. Specifically, while the original TIMIT had a large number of speakers (600) read a relatively small number (10) of sentences each, Global TIMIT reduces the recruiting effort needed to acquire an equivalent volume of data by eliciting 120 sentences from 50 speakers. In addition, Global TIMIT does not require that phonetically rich or balanced or representative sentences are created for each language as were the Harvard Sentences read in the original TIMIT. Rather, corpus designers typically select sentences of reasonable length from existing open sources (e.g. Wikipedia, lists of proverbs, etc.) that they subsequently filter to remove sentences that contain foreign or unusual words or would be otherwise difficult to read fluently. Such filtering can include selecting for phonetic

balance or representativeness from this naturally occurring text.

The Linguistic Data Consortium and Xi'an Jiaotong University used the Global TIMIT methodology to create the Mandarin Chinese - Guanzhong Dialect corpus (Jiang et al., 2020) consisting of ~five hours of read speech and transcripts. 3220 sentences were selected from the Chinese Gigaword Fifth Edition (LDC2011T13) corpus of news text. 25 females and 25 males who spoke the Guanzhong dialect each read 120 sentences where 20 sentences were read by all speakers, 40 sentences were read by 10 speakers, and 60 sentences were read by one speaker.

The corpus was recorded in a quiet room at Xi'an Jiaotong University, Xi'an, China. The collection yielded 5999 utterances (one missing). Each utterance appears in its own audio file and is accompanied by time aligned transcripts for each word, phone and tone segment as well as Praat TextGrids.

4. Collection Protocol: Object Naming

The next (i.e. current) data collection, changed a number of parameters to address details of the language that were not the focus of our previous effort. As noted above, two of the most striking differences between the Guanzhong dialect and Standard Mandarin are in the tonal system and lexicon. Speech elicitation through the reading of previously written sentences would be expected to produce data suitable for the study of phonetic and phonological differences, at least to the extent that native speakers produce such differences in read speech. However, in order to gather data on variation in the lexicon, we would need a task in which the contributors were freer to make their own lexical choices. After considering both picture description and object naming, we settled on the latter which would also yield also a picture/speech database that can be used to build visually-grounded speech models (Scholten, Merckx, Scharenborg, 2021). Initial targets were 200 objects named by 20 speakers each.

We then considered gathering images to match the items in the Communicative Development Inventories adapted for Mandarin (Tardif and Fletcher 2008) but decided to use MultiPic “a standardized set of 750 drawings with norms for six European languages” (Duñabeitia et al. 2018) and ongoing effort collecting written namings in additional languages including Chinese (Duñabeitia, p.c.). MutiPic’s set of consistently composed colored line drawings saved us the effort of identifying images to use in the object naming but, of course, required adaptation for Guanzhong dialect speakers.

4.1 LanguageARC for Collection and Annotation

The requirements for a platform to collect judgements concerning the appropriateness of the MultiPic images – and eventually for the speech samples themselves – were several. The platform would need to perform adequately in multiple locations in the USA, Netherlands and China where members of the team were located. It would need to be able to collect multiple choice, unconstrained text and speech and accept and display text in Roman and Chinese characters. It should also allow project managers to restrict access to different tasks within the project to different sets

of users. Ease of use for project managers and contributors, and the ability to rapidly prototype tasks were important considerations for this somewhat innovative effort. Finally, as the project, from initial conception to the release of the data, took place during the COVID-19 epoch, we could not ask participants to visit a carefully instrumented lab. Instead the platform needed to support entirely remote operation. After considering some alternatives, we prototyped the collection of multiple choice judgements and unconstrained text and speech in the LanguageARC citizen science portal (Cieri and Fiumara, 2020) where performance was sufficient to proceed.

4.2 Image Selection

As above the MultiPic corpus offers us a large number of images that had been normed across multiple European languages. However, it had not yet been adapted for potential contributors living in China. We had anticipated that images of the following types might prove confusing to some speakers:

1. complex images where it's not clear what should be named (#22, a campsite with multiple tents, trees, hills)
2. images specific to a time, place or domain that might not be generally known (#051, a gallows)
3. images that use a 'visual language of deixis' that might not be obvious to users when they first encounter them but may be learned over time (#3, a step on a staircase is shaded) leading to variation across the collection

To select images suitable for presentation in an object naming task to a large number of native Chinese speakers, we first presented each image to a single native speaker of Chinese, living in China, to judge whether others would be able to name the object consistently. Figure 2 shows the LanguageARC task used to present images to and collect judgements from the annotator.

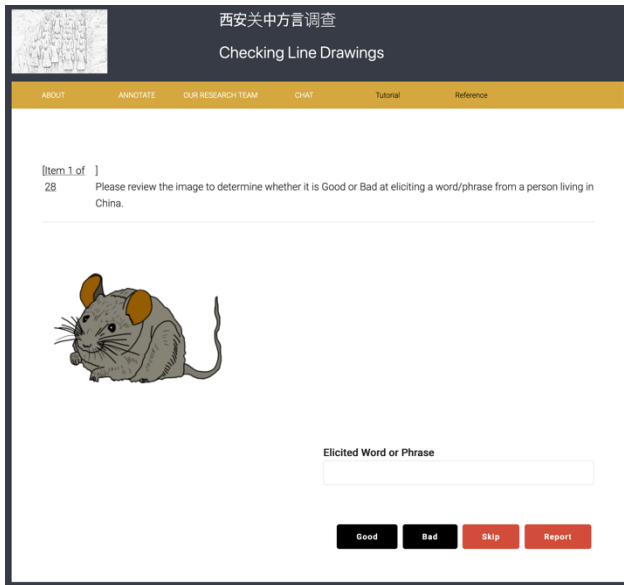


Figure 2: Selecting Appropriate Images

The annotator provided both a judgement as to appropriateness and one or more written labels for the image. A member of the project team, with similar

linguistic and cultural background, plus an understanding of our research goals, using the LanguageARC task displayed in Figure 1, reviewed those cases where the annotator expressed uncertainty and made final decisions as to whether the images would be included.

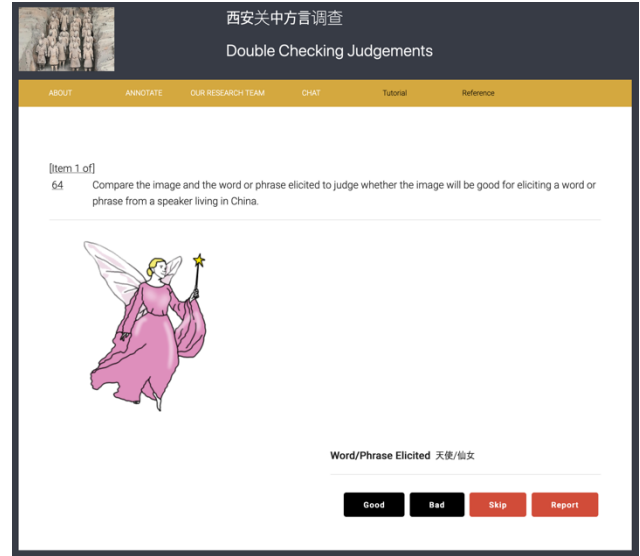


Figure 1: Double Checking Uncertain Judgements

After these two passes, we had identified 622 of the 750 MultiPic images that we believed would easily elicit object namings from Chinese contributors. Most of the excluded images were complex (103, art gallery), tightly connected to Western culture (144, toreador), religion (742, pope) or mythology (184, witch and cauldron) or perhaps just embarrassing (65, buttocks). As our goal was simply to identify many images for the elicitation we did not explore the reasons for exclusion in great detail.

5. Recruitment and incentives

As with the Mandarin Chinese - Guanzhong Dialect corpus, the present dataset was also collected by Juhong Zhan, Yue Jiang and their students from Xi'an Jiaotong University.

5.1 Recruitment

We recruited speakers of the Guanzhong dialect who had grown up in cities, counties, towns and countryside in the Guanzhong region, mainly Xi'an, Baoji, Xianyang, Weinan, Tongchuan and Yangling, etc., by posting recruitment notices on WeChat, a popular social media platform in China. For the sake of organization and management, we mainly targeted students from Xi'an Jiaotong University. The speakers needed to meet the following criteria

1. being born and growing up in Guanzhong region;
2. speaking Guanzhong dialect on a daily basis;
3. understanding English and being able to use a computer.

Participants were undergraduate students from Xi'an Jiaotong University and community members. They speak Guanzhong dialect with their families and town fellows, and Mandarin Chinese with teachers, classmates and friends who cannot understand and speak Guanzhong Dialect. They are able to freely switch between dialect and standard Mandarin.

The first task for the native Guanzhong dialect speakers was to record themselves naming the objects presented on a computer screen while speaking in their dialect (see Section 6). Each speaker was to identify and name each of 622 images using the Guanzhong dialect presented by the LanguageARC citizen science platform. We added the participants into a TenCent QQ group chat and sent them the project website and a video manual. All the speakers' questions and problems in the process of recording were answered online.

Once the recording was complete, the second task for the native speakers was quality control of the Guanzhong dialect (see §7). After automated and human effort at LDC to eliminate recordings containing digital artifacts, signal and noise problems, etc., native Guanzhong speakers were asked to judge whether the naming of the image was correct, whether the speaker was naming and not talking about the image, etc. In this second round, a total of 26 speakers participated, including some volunteer participants from the first round, and some newly recruited ones, to identify invalid recordings from those finished in the first, recording round. A tutorial video was posted with instructions on how to tag invalid recordings and the criteria for invalid recordings.

We recruited 48 participants for the recording task, 21 of whom participated in the quality control task. 18 male speakers and 30 females, mostly aged between 18-24, with only 3 over 30 participated in the recording task. Of those, 7 males and 14 females, mostly between the ages of 19-23, with one over 30 participated in the quality control task.

5.2 Incentives

A local language documentation task naturally presents several potential incentives to participants who may be attracted by their intellectual interest in language generally or in the specific dialect or by local pride and the opportunity to contribute to language preservation or by the task itself or the opportunity to work with others of like mind. Because the project was under time pressure to prepare a resource for possible use in a joint research project, we opted to augment these natural incentives of citizen science efforts.

Participants in the first task were awarded 100 yuan each, and since the recruited members were mostly students, two extra points were given to their daily performance as a reward for participation. In the quality control task, each participant tagging 1500 HITs received a standard reward of 300 yuan, with more pay for more work.

At the end of the assignment, we thanked the participants and paid them for their efforts one by one through online payment.

6. Collection

During collection we presented images to participants, one at a time, in random order and asked them to record themselves naming the items using the LanguageARC task picture in Figure 3. Participants could listen to their recording before submitting and make a new recording if they felt it necessary. Each naming took a few seconds to accomplish. Contributors could proceed at their own pace, skip items, leave the task and return as they wished. The

task included a tutorial in standard Chinese on making high quality audio recordings. The speech collection ran from February through May 2021. As this effort was framed as a citizen science project in which maximizing participation is a goal, we did not require participants to use any specific computer hardware, browser or microphone.



Figure 3: Recording Object Namings

Throughout the collection (and QC) phases researchers at Xi'an Jiaotong University worked closely with LDC to discuss any problems that might arise. Based on our experience in the pilot task and reviews of early recordings, we identified and constantly reminded the participants of these important issues:

1. speaking too close to the microphone and/or too lead can lead to clipping;
2. low signal can result from the speaker being too far from the microphone, turning away from the microphone or speaking too softly;
3. environmental noise could reduce the value of the recordings;
4. the frequency of digital artifacts was significantly higher for some participants or environments; some were ask to try a different local or device;
5. forgetting to click "Submit" after recording could result in long recordings with little useable data;
6. clicking "Skip" too frequently cold skew the distribution of the data; participants who did so were encouraged to consult the team leader;
7. recordings could be reviewed and new recordings made as needed; participants were reminded to do this if they were uncertain about their naming or the system performance.

7. Quality Control

Initial review of the audio collected via LanguageARC, from test subjects in the US and China, revealed that quality was generally quite good with high SNR. However, as the number and diversity of speakers, systems, locations and clips grew, notwithstanding the care taken during collection, quality naturally varied according to

participants' hardware and network capability, physical environments and behaviors.

Although LanguageARC relied upon an open source audio collection framework intended for use over the Internet, we experienced two types of problems: long sequences of samples with a value of 0 (nulls), presumably due to aggressive noise suppression, and stark discontinuities in the waveform, presumably due to packet loss, replication or interpolation. Although we have since adjusted the framework's parameters within LanguageARC to mitigate these problems, there were nonetheless audio samples that needed to be set aside.

At the scale the project worked, it seemed improbable that we could quality control all clips with human effort. Subsequently, we also learned that some recording problems were difficult to detect just by listening. For example, one of the commonest error patterns was a string of 127 null samples which, at the sampling rate used by the recording framework was just 2 milliseconds of silence and was often missed by human listeners. In addition, because collection had gone much better than expected (more speakers naming more objects) we could afford to set aside audio clips that revealed any undesirable properties and still have a suitable corpus. With this in mind we created a cascade of automatic and human QC processes that was intended to identify and set aside all recordings with digital artifacts and present the remainder for more careful review.

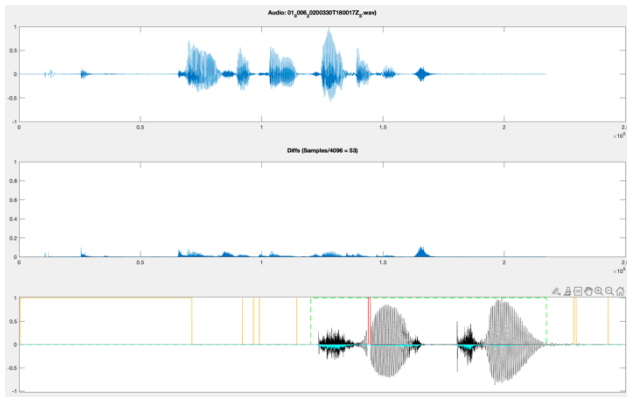


Figure 4: Automated Detection of Digitization Artefacts

7.1 Automated

To assure rapid QC of the very large number of clips collected we developed customized automated detectors in MatLab to identify any cases in which the audio contained sequences of zero-valued samples or large discontinuities in adjacent sample values during speech segments. These automatic detectors were correlated with human judgments of the same types of problems and adjusted until the automated methods were at least as sensitive as human ears. At that point, the automated detectors were used to preprocess the entire corpus. Any audio clip that suffered from either problem type was set aside and the remaining clips were subjected to further human QC. Figure 4 shows an interface used in developing the automated detectors. The panels, from top to bottom, contain: the waveform; a plot of the difference between adjacent sample values where large discontinuities appear as spikes (none apparent in this example); and a zoomed in display of a waveform showing a speech region (green box) and null sequences

(orange and red boxes, the latter indicating null strings during speech).

7.2 Human Quality Control

In the second round, the criteria for invalid recordings were set as follows: recordings being incomplete or distorted; speech being too low and soft; noise being too loud; recordings in a dialect other than Guanzhong; or a naming mismatch to the image presented.

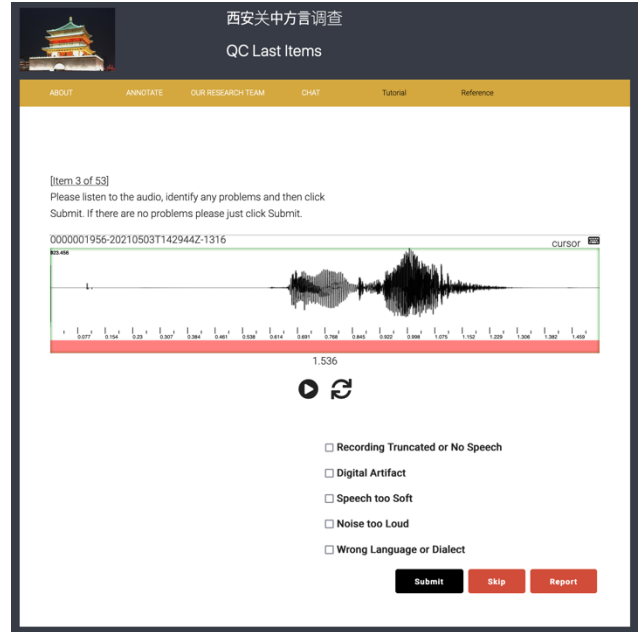


Figure 5: Human Quality Control of Audio Recordings

8. Resulting Data

The object naming task yielded 34,729 audio recordings in total. After automated quality control excluded clips for sequences of nulls or large discontinuities in the wave form and human quality control excluded clips that were truncated or still sounded as if they were marred by digital artifacts, 25,972 remained. Any of these annotated for the other problems: speech too soft, background noise too loud or contained lexical items in the wrong language/dialect were retained but appropriated marked in the metadata tables.

The corpus is organized into 622 directories according to the image presented. Each directory contains on average 42 recordings of namings of the object in the image (min=7, max=54, std=4.6) sampled at 16kHz, 16bit, single channel, WAV files. Files are named to indicate the image presented, a userID and the date and time of the recording. An accompanying table indicates, for each file, whether any human annotator indicated high noise, low signal or that the items is not in the target dialect. The table also provides the file size, duration of the audio file, duration of the speech portion and pseudo-SNR.

The corpus was originally released by LDC in September 2021 to participants in the Frontiers in AI Research Topic in *AI and Low-resource Speech Sciences* and will shortly be released to the entire research community.

9. Acknowledgements

The authors would like to express their gratitude to the US National Science Foundation (including CRI grant 1730377 which supports the NIEUW project described above); the Netherlands Organization for Scientific Research (NWO-VI.C.181.040) that funded related research on picture naming and data collection in understudied languages; the University of Pennsylvania Office of the Vice Provost for Global Initiatives (PennGlobal) through its Penn China Research and Engagement Fund, and the University of Pennsylvania School of Arts and Sciences through its Global Engagement Fund that supported the creation of the Documenting Xi'an Guanzhong - Object Naming corpus and the Mandarin Chinese - Guanzhong Dialect corpus.

10. Bibliographical References

- Bai, D. (1954). A Report on Guanzhong Dialect, Beijing, Chinese Academy of Sciences Press (in Chinese).
- Chanchaochai, N., Cieri, C., Debrah, J., Ding, H., Jiang, Y., Liao, S., Liberman, M., Wright, J., Yuan, J., Zhan, J., Zhan, Y. (2018) GlobalTIMIT: Acoustic-Phonetic Datasets for the World's Languages, *Proceedings of Interspeech*.
- Cieri, C., Fiumara, J. (2020) LanguageARC – a tutorial LREC 2020: 12th Edition of the Language Resources and Evaluation Conference, CLLRD Workshop: Citizen Linguistics in Language Resource Development
- Cieri, C., Fiumara, J., and Wright, J. (2021). Using Games to Augment Corpora for Language Recognition and Confusability, *Proc. 22nd Annual Conference of the International Speech Communication Association (Interspeech)*, August 30-September 3.
- Duñabeitia, J.A., Crepaldi, D., Meyer, A. S., New, B., Pliatsikas, C., Smolka, E., and Brysbaert, M. (2018). MultiPic: A standardized set of 750 drawings with norms for six European languages. *Quarterly Journal of Experimental Psychology*. 71(4):808-816.
- Fiumara, J., Cieri, C., Wright, J., and Liberman, M. (2020). LanguageARC: Developing Language Resources Through Citizen Linguistics in *Proc. 12th Edition of the Language Resources and Evaluation Conference (LREC)*. CLLRD Workshop: Citizen Linguistics in Language Resource Development. Marseille, May 11-16.
- Joglekar, A, Seyed, O. S., Chandra-Shekar, M., Cieri, C., and Hansen, J.H.L. (2021). Fearless Steps Challenge Phase-3 (FSC P3): Advancing SLT for Unseen Channel and Mission Data Across NASA Apollo Audio in *Proc. 22nd Annual Conference of the International Speech Communication Association (Interspeech)*, August 30-September 3.
- Li, R. (1989) Language Atlas of China, Longman.
- Li Min. (2014). Study on Major Differences between Typical words in Guanzhong Dialect and Mandarin. *Literature Education* (01), 124-125.
- Liu, M., Chen, Y., Schiller, N. O. (2020). Tonal mapping of Xi'an Mandarin and Standard Chinese. *The Journal of the Acoustical Society of America*, 147(4), 2803. <https://doi.org/10.1121/10.0000993>.
- Lu Tuanhua. (2010). Comparison of the Phonetic Characteristics between Guanzhong Dialect and Mandarin. *KAOSHI ZHOUKAN* (09), 23-24.

- Scholten, S., Merkx, D., Scharenborg, O., (2021) Learning to Recognise Words using Visually Grounded Speech, *IEEE International Symposium on Circuits and Systems (ISCAS)*.
- Sun Lixin. (2021). Complementary Discussions on Personal Pronouns in Guanzhong Dialect. *Journal of Gansu Normal Colleges* (01), 1-5.
- Tardif, T. and Fletcher, P. *Chinese Communicative Development Inventories: User's Guide and Manual* (2008), Peking University Medical Press, Beijing, China.
- Wang Wenya. (2015). Study on Phonetic and Lexical Features and History of Guanzhong dialect in Shaanxi Province. *China Juveniles* (22), 169-170.
- Wang Yuding. (1995). Studies On Types of Sound Changes in the Guanzhong Dialect. *Journal of Yan'an University (Social Sciences Edition)*.
- Xing Xiangdong. (2014). The Key Investigations and Researches on Dialects in Northwest Regions—Mainly on Gansu, Ningxia, Qinghai, and Xinjiang Provinces. *Journal of Tsinghua University (Philosophy and Social Sciences)* (05), 122-134+178. Doi: 10.13613/j.cnki.qhdz.002267.
- Zhang, W. J. (2005). Drifting and Competition: The Changes of the phonological structure of Guanzhong dialect. Shaanxi People's Publishing House.
- Zhao Nana. (2020). A Research Review on Guanzhong Dialect Vocabulary. *Journal of Jilin TV & Radio University* (04), 140-142.
- Zhao Yonggang. (2017). Chinese Tone Sandhi and the Regional Differences of Citation Tones of Shaanxi Dialect. *Foreign Language and Literature Research*. 3(4), 12.

11. Language Resource References

- Cieri, C., Zhan, J., Jiang, Y., Liberman, M., Yuan, J., Chen, Y., Scharenborg, O. (2021) Documenting Xi'an Guanzhong - Object Naming LDC2021E14. Web Download. Philadelphia: Linguistic Data Consortium.,
- Garofolo, J., Lamel, L. Fisher, W., Fiscus, J., Pallett, D., Dahlgren, N., Zue, V. (1993) TIMIT Acoustic-Phonetic Continuous Speech Corpus LDC93S1. Web Download. Philadelphia: Linguistic Data Consortium.
- Jiang, Y., Zhan, J., Han, H., Xu, Z., Zhou, H., Yuan, J., Liberman, M. (2020) Global TIMIT Mandarin Chinese-Guanzhong Dialect LDC2020S12. Web Download. Philadelphia: Linguistic Data Consortium.