

URN: urn:nbn:de:kobv:b4-opus-24433

BRYAN JURISH,
Canonicalizing the Deutsches Textarchiv,

in:

*Perspektiven einer corpusbasierten historischen Linguistik und Philologie.
Internationale Tagung des Akademienvorhabens „Altägyptisches Wörter-
buch“ an der Berlin-Brandenburgischen Akademie der Wissenschaften,
12. – 13. Dezember 2011, herausgegeben von Ingelore Hafemann,
Berlin 2013, S. 235-244.*

BERLIN-BRANDENBURGISCHE AKADEMIE DER WISSENSCHAFTEN

Thesaurus Linguae Aegyptiae 4

Perspektiven einer corpusbasierten historischen Linguistik und
Philologie. Internationale Tagung des Akademienvorhabens
„Altägyptisches Wörterbuch“ an der Berlin-Brandenburgischen
Akademie der Wissenschaften, 12. – 13. Dezember 2011

herausgegeben von Ingelore Hafemann

BERLIN-BRANDENBURGISCHE AKADEMIE DER WISSENSCHAFTEN

Thesaurus Linguae Aegyptiae

4

BERLIN 2013

BERLIN-BRANDENBURGISCHE AKADEMIE DER WISSENSCHAFTEN

Perspektiven einer corpusbasierten historischen Linguistik
und Philologie

Internationale Tagung des Akademienvorhabens „Altägyptisches
Wörterbuch“ an der Berlin-Brandenburgischen Akademie der
Wissenschaften, 12. – 13. Dezember 2011

herausgegeben von Ingelore Hafemann

BERLIN

2013



Dieser Band wurde durch die gemeinsame Wissenschaftskonferenz im Akademienprogramm mit Mitteln des Bundes (Bundesministerium für Bildung und Forschung) und des Landes Berlin (Senatsverwaltung für Wirtschaft, Technologie und Forschung) gefördert

Die Publikation unterliegt folgender Creative-Commons-Lizenz:
„Namensnennung – Keine kommerzielle Nutzung – Weitergabe unter gleichen Bedingungen 3.0 Deutschland“

<http://creativecommons.org/licenses/by-nc-sa/3.0/de/>



URN: urn:nbn:de:kobv:b4-opus-24310

INHALTSVERZEICHNIS

VORWORT	7
GREGORY CRANE & ALISON BABEU Global Editions and the Dialogue among Civilizations	11
HISTORISCHE CORPUS-PROJEKTE – SYNCHRON UND DIACHRON	
STÉPHANE POLIS & JEAN WINAND The Ramses project. Methodology and practices in the annotation of Late Egyptian Texts	81
SERGE ROSMORDUC The Ramses project in perspective. Managing evolving linguistic data	109
DIETER KURTH Das Edfu-Projekt. Ziel, Methode und Verarbeitung der lexikographischen Ergebnisse	121
INGELORE HAFEMANN & PETER DILS Der Thesaurus Linguae Aegyptiae – Konzepte und Perspektiven	127
GÜNTER VITTMANN Zur Arbeit an der Demotischen Textdatenbank: Textauswahl	145
GERNOT WILHELM Das Hethitologie Portal Mainz	155
JOST GIPPERT The TITUS Project. 25 years of corpus building in ancient languages	169
KURT GÄRTNER & RALF PLATE Die Doppelfunktion des digitalen Textarchivs als Wörterbuchbasis und als Komponente der Online-Publikation. Am Beispiel des Mittelhochdeutschen Wörterbuchs	193
HANS-CHRISTIAN SCHMITZ, BERNHARD SCHRÖDER & KLAUS-PETER WEGERA Das Bonner Frühneuhochdeutsch-Korpus und das Referenzkorpus ‚Frühneuhochdeutsch‘	205

ALEXANDER GEYKEN Wege zu einem historischen Referenzkorpus des Deutschen: das Projekt Deutsches Textarchiv	221
BRYAN JURISH Canonicalizing the Deutsches Textarchiv	235
WORTGESCHICHTE - TEXTGESCHICHTE - SPRACHGESCHICHTE: TRADITION UND INNOVATION BEI DER TEXTPRODUKTION	
FRANK FEDER & SIMON D. SCHWEITZER Auf dem Weg zu einem integrierten Lexikon des Ägyptisch- Koptischen	245
FRIEDHELM HOFFMANN Die Demotische Wortliste – virtuell erweitert	263
GÜNTER VITTMANN Kursivhieratische Texte aus sprachlicher und onomastischer Sicht	269
MATHEW ALMOND, JOOST HAGEN, KATRIN JOHN, TONIO SEBASTIAN RICHTER & VINCENT WALTER Kontaktinduzierter Sprachwandel des Ägyptisch-Koptischen: Lehnwort-Lexikographie im Projekt Database and Dictionary of Greek Loanwords in Coptic (DDGLC)	283
THOMAS GLONING Historischer Wortgebrauch und Themengeschichte. Grundfragen, Corpora, Dokumentationsformen	317
LOUISE GESTERMANN Die altägyptischen Sargtexte in diachroner Überlieferung	371
THOMAS STÄDTLER Überlegungen zu Textsorte und Diskurstradition bei der Beschreibung von Textcorpora und ihr Bezug zur lexikographischen Forschung	385

VORWORT

Die internationale Tagung „Perspektiven einer corpusbasierten historischen Linguistik und Philologie“ vom 12. – 13. Dezember 2011 am Akademienvorhaben „Altägyptisches Wörterbuch“ der Berlin-Brandenburgischen Akademie der Wissenschaften (BBAW) war dem Thema des Aufbaus und der Nutzungsperspektiven elektronischer Textcorpora und Wörterbücher in den historischen Sprachen gewidmet. Die Teilnehmer, Vertreter der Ägyptologie, der Hethitologie, Indogermanistik sowie Referenten aus der historischen Lexikographie des Mittel- und Frühneuhochdeutschen und des Altfranzösischen diskutierten vor allem über die Veränderungen, die mit dem Einsatz elektronischer Erfassungs- und Verarbeitungsprozeduren einhergehen. Vertreter der Computerlinguistik vom „Zentrum Sprache“ der BBAW wurden in die Diskussionen einbezogen. Dort beschäftigt man sich seit Jahren mit dem Aufbau großer elektronischer Textcorpora (DWDS), darunter auch solcher, die historische Texte (DTA) für die elektronische Nutzung ermöglichen.

Die größte Herausforderung dieser neuen elektronischen Corpora und Wörterbücher ist es, sowohl den Methoden und damit den wissenschaftlichen Ansprüchen der traditionellen Philologie und Lexikographie unbedingt verpflichtet zu bleiben als auch neue Gebiete wie die Corpus- und Computerlinguistik für die historischen Sprachen zu öffnen. Die Teilnehmer haben gemeinsam und disziplinenübergreifend die Möglichkeiten und Grenzen der Datenerfassung, ihrer Präsentation und den Nutzen neuer Auswertungsprozeduren diskutiert.

Unter dem ersten Thema „Historische Corpusprojekte – synchron und diachron“ wurden elektronische Corpora vorgestellt und ein intensiver Austausch darüber geführt, welche Datenstrukturen die linguistischen Inhalte in adäquater Weise abbilden. Wichtig war die Frage, auf welche Resonanz diese elektronischen Corpora bei den Nutzern gestoßen sind und welche Erwartungen und Anforderungen aus den verschiedenen Fachdisziplinen an die Projekte herangetragen werden. Der Austausch über Nutzungsperspektiven elektronischer Corpora schloss auch die Diskussion über die Erarbeitung projektübergreifend einsetzbarer Standards der Codierung und Strukturierung historischer Textdaten mit ein. Hinsichtlich einer mittel- und langfristigen Nutzbarkeit sowie einer langfristigen Datensicherheit stehen solche Fragen zunehmend im Focus und einige aktuelle Initiativen dazu wurden vorgestellt. Spezielle technische Aspekte

elektronischer Datenerfassung und automatischer Analyse- und Speicherungsverfahren elektronischer Textdaten konnten am letzten Tag als ein Themenschwerpunkt mit den Programmierern diskutiert werden.

Ein zweiter Schwerpunkt waren konkrete Fragestellungen aus der historischen Lexikographie und diachronen Textanalyse. Für das Ägyptische ist der diachrone Ansatz auf Grund der über vier-tausendjährigen Textüberlieferung von großer Relevanz. Themen wie historischer und/oder textgattungsspezifischer Wortgebrauch, die Erarbeitung diachroner Wortlisten und Aspekte des kontaktindizierten Sprachwandels konnten disziplinübergreifend zwischen den Ägyptologen und den Kollegen der historischen Lexikographie des Mittel- und Frühneuhochdeutschen und des Altfranzösischen behandelt werden.

Mit dem Abendreferenten Gregory Crane, dem Begründer der „Perseus Digital Library“, wurde ein breites Publikum angesprochen. In seinem Vortrag hat er noch einmal die hohe Relevanz und die neuen Möglichkeiten der Einbeziehung zahlreicher Wissenschaftler und einer interessierten Öffentlichkeit in die Projektarbeit demonstriert, die das Internet auf völlig neue Weise eröffnet hat. Die Herausgeberin ist sehr froh, seinen programmatischen Beitrag zu diesem Thema, dessen schriftliche Form er gemeinsam mit Alison Babeu erarbeitet hat, ebenfalls in diesem Band präsentieren zu können.

Wir danken der Berlin-Brandenburgischen Akademie der Wissenschaften für die umfassende Unterstützung unserer Projektarbeit und ganz speziell der Vorbereitung dieser Konferenz sowie der Möglichkeit, die Akten auf dem E-Doc-Server der Akademie veröffentlichen zu können.

Der Hermann und Elise geborene Heckmann Wentzel-Stiftung sei hiermit ausdrücklich für die unbürokratische und großzügige finanzielle Unterstützung dieser erfolgreichen Tagung gedankt.

Das Akademienvorhaben „Altägyptisches Wörterbuch“ konnte sich als aktives Mitglied des Weiteren auf das „Zentrum Grundlagenforschung Alte Welt“ stützen, dem alle altertumswissenschaftlichen Vorhaben der BBAW angehören. Dem Zentrum ist es zu danken, dass der Abendvortrag von Gregory Crane einem breiteren Publikum dargeboten werden konnte.

Allen Autoren dankt die Herausgeberin für ihre anregenden Diskussionen und die qualitätvollen Beiträge in diesem Band.

Auf eine Gesamtbibliographie wurde verzichtet und die Abkürzungen der in den ägyptologischen Beiträgen erwähnten Zeitschriften und Reihen folgen dem Lexikon der Ägyptologie, herausgegeben von Wolfgang Helck und Wolfhart Westendorf, Band VII: Nachträge, Korrekturen, Indices, Wiesbaden 1992, XIV-XIX.

Ganz besonders sei schließlich Frau Angela Böhme für die gewissenhafte redaktionelle Bearbeitung der Manuskripte gedankt sowie Dr. Simon Schweitzer für seine Hilfe beim Erstellen des Layouts.

Berlin, Mai 2013

Ingelore Hafemann

CANONICALIZING THE DEUTSCHES TEXTARCHIV

BRYAN JURISH

1. Introduction

Virtually all conventional text-based natural language processing techniques – from traditional information retrieval systems to full-fledged parsers – require reference to a fixed lexicon accessed by surface form, typically trained from or constructed for synchronic input text adhering strictly to contemporary orthographic conventions. Unconventional input such as historical text which violates these conventions therefore presents difficulties for any such system due to lexical variants present in the input but missing from the application lexicon.

Traditional approaches to the problems arising from an attempt to incorporate historical text into such a system rely on the use of additional specialized (often application-specific) lexical resources to explicitly encode known historical variants. Such specialized lexica are not only costly and time-consuming to create, but also – in their simplest form of static finite word lists – necessarily incomplete in the case of a morphologically productive language like German, since a simple finite lexicon cannot account for highly productive morphological processes such as nominal composition [KEMPEN *et al.* (2006)].

To facilitate the extension of synchronically-oriented natural language processing techniques to historical text while minimizing the need for specialized lexical resources, one may first attempt an automatic canonicalization of the input text. Canonicalization approaches treat orthographic variation phenomena in historical text as instances of an error-correction problem, seeking to map each (unknown) word of the input text to one or more extant *canonical cognates*: synchronically active types which preserve both the root and morphosyntactic features of the associated historical form(s). To the extent that the canonicalization was successful, application-specific processing can then proceed normally using the returned canonical forms as input, without any need for additional modifications to the application lexicon.

This paper provides an informal overview of the various canonicalization techniques currently employed by the *Deutsches*

*Textarchiv*¹ (DTA; [GEYKEN & KLEIN (2010)]) project at the Berlin-Brandenburg Academy of Sciences and Humanities to prepare a corpus of historical German text for part-of-speech tagging, lemmatization, and integration into a robust online information retrieval system. For more details on the methods employed, the interested reader is referred to [JURISH (2012)].

2. Canonicalization Techniques

It is useful to distinguish between *type-wise* and *token-wise* canonicalization techniques. Type-wise canonicalization techniques are those which process each input word in isolation, independently of its surrounding context, and are fully specified by a binary *conflation relation*² over surface strings. Token-wise canonicalization techniques on the other hand make use of the context in which a given instance of a word occurs when determining the optimal canonical cognate, and can thus better account for ambiguities in the mapping from historical to contemporary forms, insofar as these can be resolved by reference to the immediate context. In the sequel, $w \sim_r v$ indicates that the words (types or tokens) w and v are related by the conflator r , and $[w]_r$ denotes the set of all words conflated with w by the conflator r . If r returns a unique (canonical) value v for each input word w , the standard notation for functions $r(w) = v$ will be used.

2.1 String Identity

String identity (*id*) is the easiest conflator to implement (no additional programming effort or resources are required) and provides a high degree of precision, “false friends” being limited to historical homographs such as the historical form *wider* when it occurs as a variant of the contemporary form *wieder* (“again”) rather than the lexically distinct contemporary homograph *wider* (“against”). Since its coverage is restricted to valid contemporary forms, string identity cannot account for any spelling variation at all, resulting in very poor recall – many relevant types will not be retrieved in response to a query in current orthography.

As an example, consider the historical form *Abſt nde*, a variant of the contemporary cognate *Abst nde* (“distances”). The conflation set $[Abſt nde]_{id} = \{Abſt nde\}$ does not contain the desired contemporary cognate, so no instances of the historical variant *Abſt nde* will be

¹ “German Text Archive”, <http://www.deutschestextarchiv.de>.

² Prototypically, every conflation relation will be a true equivalence relation.

retrieved via string identity for a query of the contemporary form *Abstände*. In the DTA canonicalization architecture, string identity is used only as a fallback conflator. Each input word is treated as its own canonical form if all other canonicalizations methods have failed, or if it passes some simple heuristic tests for detecting “uncanonicalizable” strings such as punctuation, abbreviations, mathematical formulæ, or foreign language material in non-latin script.

2.2 Transliteration

A slightly less naïve family of conflation methods are those which employ a simple deterministic transliteration function to replace input characters which do not occur in contemporary orthography with extant equivalents. A transliteration conflator is defined in terms of a character transliteration function *xlit* which maps each possible input character to a (possibly empty) output string over the contemporary alphabet, the concatenation of which yields the candidate canonical form for the input word.

In the case of historical German, deterministic transliteration is especially useful for its ability to account for typographical phenomena, e.g. by mapping ‘ſ’ (long ‘s’, as commonly appeared in texts typeset in *Fraktur*) to a conventional round ‘s’, and mapping superscript ‘e’ to the conventional *Umlaut* diacritic “”, as in the transliteration $\text{xlit}(\text{Abſt\ddot{a}nde}) = \text{Abstände}$ (“distances”). Given this transliteration, a query for the contemporary form *Abstände* will successfully retrieve all instances of the historical form *Abſt\ddot{a}nde*.

The DTA canonicalization cascade uses a fast conservative transliteration function based on the `Text::Unidecode` Perl module.³ Despite its efficiency, and although it outdoes even string identity in terms of its precision, deterministic transliteration suffers from its inability to account for spelling variation phenomena involving extant characters such as the *th/t* and *ey/ei* allographs common in historical German. As an example, consider an instance of the historical form *Theyl* corresponding to the contemporary cognate *Teil* (“part”). Both historical and contemporary forms will be transliterated to themselves, since both strings contain only extant characters, but the historical form will not be retrieved by a query for the contemporary form, since their transliterations are distinct: $\text{xlit}(\text{Teil}) = \text{Teil} \neq \text{Theyl} = \text{xlit}(\text{Theyl})$.

³ <http://search.cpan.org/~sburke/Text-Unidecode-0.04/>.

2.3 *Phonetization*

A more powerful family of conflation methods is based on the dual intuitions that graphemic forms in historical text were constructed to reflect phonetic forms⁴ and that the phonetic system of the target language is diachronically more stable than its graphematic system. A phonetic conflator maps each input word w to a unique phonetic form $\text{pho}(w)$ by means of a computable function pho , conflating those strings which share a common phonetic form. The phonetic conversion module used in the DTA was adapted from the phonetization rule-set distributed with the IMS German Festival package [MÖHLER *et al.* (2001)], a German language module for the Festival text-to-speech system [BLACK & TAYLOR (1997)], and compiled as a finite-state transducer.⁵

Phonetic conflation offers a substantial improvement in recall over conservative methods such as transliteration or string identity: variation phenomena such as the *th/t* and *ey/ei* allographs mentioned above are correctly captured by the phonetization transducer: $\text{pho}(\textit{Theyl}) = [\text{tail}] = \text{pho}(\textit{Teil})$, which implies that all instances of the historical form *Theyl* will be retrieved in response to a query of the contemporary form *Teil*. Unfortunately, these improvements come at the expense of precision: in particular, many high-frequency types are misconflated by the simplified phonetization rule-set, including **in* ~ *ihn* (“in” ~ “him”) and **wider* ~ *wieder* (“against” ~ “again”). While such high-frequency cases can easily be dealt with by a small exception lexicon (cf. section 2.6), the underlying tendency of strict phonetic conflation either to over- or to under-generalize – depending on the granularity of the phonetization function – is likely to remain, expressing itself in information retrieval tasks as reduced precision or reduced recall, respectively.

2.4 *Rewrite Transduction*

Despite its comparatively high recall, the phonetic conflator fails to relate unknown historical forms with any extant equivalent whenever the graphemic variation leads to non-identity of the respective phonetic forms (e.g. $\text{pho}(\textit{umb}) = [\text{?ump}] \neq [\text{?um}] = \text{pho}(\textit{um})$ for the historical variant *umb* of the preposition *um* (“around”)),

⁴ [KELLER (1978)] codifies this intuition as the imperative “write as you speak” governing historical spelling conventions.

⁵ In the absence of a language-specific phonetization function, a general-purpose phonetic digest algorithm such as SOUNDEX [RUSSELL (1918)] may be employed instead [ROBERTSON & WILLETT (1993)].

suggesting that recall might be further improved by relaxing the strict identity criterion implicit in the definition of the phonetic conflator. A conflation technique which fulfills both of the above desiderata is *rewrite transduction*,⁶ which can be understood as a generalization of the well-known *string edit distance* [DAMERAU (1964), LEVENSHTAIN (1966)].

A rewrite conflator (rw) is defined in terms of a *target lexicon* of contemporary forms and a weighted *error model* [KERNIGHAN *et al.* (1990), BRILL & MOORE (2000)] which associates each known pattern of diachronic variation with a non-negative weight or “distance”. The conflation set $[w]_{rw,k}$ is computed as the set of k nearest neighbors of the input word w which are themselves members of the target lexicon. Importantly, such a rewrite conflation set can be computed even in the presence of an infinite target lexicon,⁷ provided that both lexicon and error model can be represented as (weighted) finite-state transducers [MOHRI (2002)].

The DTA canonicalization architecture uses a finite-state rewrite cascade whose error model was compiled from a set of manually constructed rules and whose target lexicon was extracted from the the high-coverage TAGH morphology system for contemporary German [GEYKEN & HANNEFORTH (2006)] to compute rewrite conflation sets containing at most only a single “best” contemporary form ($k=1$). Although this rewrite cascade does indeed improve both precision and recall with respect to the phonetic conflator, these improvements are of comparatively small magnitude, precision in particular remaining well below the level of conservative conflators such as naïve string identity or transliteration, due largely to interference from “false friends” such as the valid contemporary compound *Rockermehl* (“rocker-flour”) for the historical variant *Rockermel* of the contemporary form *Rockärmel* (“coat-sleeve”).

2.5 Hidden Markov Model Disambiguation

Systematic evaluations of the type-wise techniques described above revealed a typical precision-recall trade-off pattern: the ultra-conservative string identity conflator – despite its near-perfect precision – shows quite poor recall, while the more ambitious high-recall

⁶ Related approaches to historical variant detection include [RAYSON *et al.* (2005), ERNST-GERLACH & FUHR (2006), GOTSCHAREK *et al.* (2009)].

⁷ E.g. as arising from morphologically productive phenomena such as German nominal composition.

conflators such as phonetic identity or rewrite transduction tend to be disappointingly imprecise. In order to recover some of the precision offered by conservative conflation techniques such as transliteration while still benefiting from the flexibility and improved recall provided by more ambitious techniques such as phonetization or rewrite transduction, the DTA canonicalization architecture makes use of a Hidden Markov Model (HMM) disambiguator which operates on the token level, using sentential context to determine a unique “best” canonical form for each input token, in a manner similar to the spelling correction technique described by [MAYS *et al.* (1991)].

Treating the conflation sets returned by all active type-wise conflators as candidate canonicalization hypotheses, the HMM disambiguator chooses an optimal sequence of token-wise unique canonical forms for each input sentence by application of the well-known Viterbi algorithm [VITERBI (1967)]. Lexical probabilities are dynamically computed as a Maxwell-Boltzmann distribution over the candidate conflations for each input word, and a static trigram model of contemporary German is used to model local syntactic and semantic context constraints. The disambiguator is thus able to resolve ambiguous conflation sets such as {*in*, *ihn*} or {*wider*, *wieder*} in a context-dependent manner.

Using a simple smoothing mechanism, the disambiguator is also able to override the decisions of the type-wise conflators by selecting a canonical form not explicitly enumerated in the target lexicon.⁸ This behavior is particularly useful for proper names, which are not exhaustively represented in the TAGH morphology system, and which were excluded entirely from the rewrite target lexicon because their presence lead to too many spurious conflations with valid historical forms, e.g. the TAGH lexical entry for the surname *Aehnlich* caused the rewrite conflator to treat all instances of the type as their own canonical forms rather than mapping them to the correct contemporary form *ähnlich* (“similar”).

2.6 Exception Lexicon

The HMM disambiguator performs very well at the token level, but its reliance on a static *n*-gram model over contemporary forms is

⁸ Technically, the possibility of selecting the input word itself as its own canonical form is implemented by allowing the identity conflator *id* to provide a candidate conflation hypothesis. In practice, the DTA canonicalization architecture uses the transliteration conflator *xlit* whenever it returns a non-empty string, and only resorts to a pure identity hypothesis when the transliterator fails.

problematic for input words whose canonical cognate was not present in the training data: in such cases, the model effectively reverts to a type-level canonicalization, choosing the most likely conflation candidate based only on criteria of word length and source conflator. Due to the conflator-dependent distance functions employed, short input words in particular are likely to be subjected to such treatment, which was designed primarily to handle low-frequency unlexicalized types such as proper names, and thus often results in a fallback identity canonicalization. Many common historical variants of high-frequency words fall into this category, usually due to (irregular) patterns of variation not captured by the type-wise conflators such as exhibited by the (strongly inflected) historical variant *frug* of the (weakly inflected) contemporary form *fragte* (“asked”). On the other hand, “false friends” in the training data can cause also spurious canonicalizations: even a single occurrence of the given name *André* in the training data is sufficient to cause the historical variant *andre* of the contemporary form *andere* (“other”) to be miscanonicalized.

To handle problematic cases such as these, the DTA canonicalization architecture incorporates a semi-automatically generated exception lexicon [JURISH *et al.* (forthcoming)] which operates on incoming word types before they are passed to the disambiguator. If the exception lexicon contains an entry for an input type, only that entry is considered by the disambiguator as a candidate canonicalization for the input word. This technique ensures that the exception lexicon entry will in fact be the canonical form chosen on the one hand, and allows the disambiguator to make use of the provided entry for context-dependent resolution of nearby items on the other.

3. Summary

Historical text presents unique challenges for typical synchronically oriented natural language processing tasks. In particular, violations of contemporary orthographic conventions are problematic for any task requiring reference to a fixed lexicon keyed by surface word type. Part-of-speech tagging, lemmatization, and information retrieval (corpus indexing & query) are all affected. Canonicalization approaches address this problem by attempting to map unknown historical variants to extant contemporary forms and deferring synchronically oriented analysis to the returned (canonical) forms.

The canonicalization techniques currently used to preprocess the *Deutsches Textarchiv* corpus of historical German were briefly described. String identity on its own does not provide an adequate solution, since it cannot account for any orthographic variation at all, but it can be useful in conjunction with additional heuristics for detecting non-lexical material. Transliteration provides an efficient and very precise canonicalization method for dealing with extinct characters such as the long ‘s’ common in historical German, but cannot account for any variation involving extant characters. More ambitious techniques such as conflation by phonetic identity or rule-based rewrite transduction are able to account for a much wider range of variation, but these improvements come at the cost of precision. Use of a Hidden Markov Model to disambiguate canonicalization hypotheses at the token level using sentential context effectively recovers much of this lost precision while still benefitting from the improved recall. Remaining systematic canonicalization errors are accounted for by a type-wise exception lexicon. The fully canonicalized corpus was subsequently tagged and lemmatized before being indexed by a robust information retrieval system which uses the canonical-lemma token-level conflation relation to implement an intuitive linguistically motivated search term expansion operator for non-expert user queries.

REFERENCES

- BLACK, A. W. & P. TAYLOR, 1997: *Festival speech synthesis system*, Technical Report HCRC/TR-83, University of Edinburgh, Centre for Speech Technology Research.
URL: <http://www.cstr.ed.ac.uk/projects/festival>.
- BRILL, E. & R. C. MOORE, 2000: An improved error model for noisy channel spelling correction, in: *Proceedings of the 38th Annual Meeting of the Association for Computational Linguistics*, 286-293.
- DAMERAU, F. J., 1964: A technique for computer detection and correction of spelling errors, in: *Communications of the ACM* 7, 171-176. DOI: 10.1145/363958.363994.
- ERNST-GERLACH, A. & N. FUHR, 2006: Generating search term variants for text collections with historic spellings, in: LALMAS, M. *et al.* (eds.), *Advances in Information Retrieval*, Lecture Notes in Computer Science 3936, Berlin, 49-60.
DOI: 10.1007/11735106_6.
- GEYKEN, A. & T. HANNEFORTH, 2006: TAGH: A complete morphology for German based on weighted finite state automata, in: *Finite State Methods and Natural Language Processing, 5th International Workshop, FSMNLP 2005, Revised Papers*, Lecture Notes in Computer Science 4002, Berlin, 55-66.
DOI: 10.1007/11780885_7.
- GEYKEN, A. & W. KLEIN, 2010: Deutsches Textarchiv, in: *Jahrbuch 2009*, Berlin-Brandenburgische Akademie der Wissenschaften, Berlin, 320-323. URL: http://edoc.bbaw.de/volltexte/2010/1515/pdf/BBAW_Jahrbuch_2009.pdf.
- GOTSCHAREK, A. *et al.*, 2009: Enabling information retrieval on historical document collections: the role of matching procedures and special lexica, in: *Proceedings of The Third Workshop on Analytics for Noisy Unstructured Text Data, AND '09*, New York, 69-76. DOI: 1568296.1568309.
- JURISH, B., 2012: Finite-State Canonicalization Techniques for Historical German, PhD thesis, Universität Potsdam.
URL: <http://opus.kobv.de/ubp/volltexte/2012/5578/>.
- JURISH, B. *et al.*, forthcoming: Constructing a canonicalized corpus of historical German by text alignment, in: BENNETT, P. *et al.* (eds.),

New Methods in Historical Corpus Linguistics, vol. 3 of *Corpus Linguistics and Interdisciplinary Perspectives on Language (CLIP)*, Tübingen.

KELLER, R. E., 1978: *The German Language*, London.

KEMPKEN, S. *et al.*, 2006: Comparison of distance measures for historical spelling variants, in: BRAMER, M. (ed.), *Artificial Intelligence in Theory and Practice*, Boston, 295–304.
DOI: 10.1007/978-0-387-34747-9_31.

KERNIGHAN, M. D. *et al.*, 1990: A spelling correction program based on a noisy channel model, in: *Proceedings COLING-1990*, vol. 2, 205–210.

LEVENSHTAIN, V. I., 1966: Binary codes capable of correcting deletions, insertions, and reversals, in: *Soviet Physics Doklady* 10, 707–710.

MAYS, E. *et al.*, 1991: Context based spelling correction, in: *Information Processing & Management* 27, 517–522.
DOI: 10.1016/0306-4573(91)90066-U.

MÖHLER, G. *et al.*, 2001: *IMS German Festival manual, version 1.2*. Institute for Natural Language Processing, University of Stuttgart.
URL: <http://www.ims.uni-stuttgart.de/phonetik/synthesis>.

MOHRI, M., 2002: Semiring frameworks and algorithms for shortest-distance problems, in: *Journal of Automata, Languages and Combinatorics* 7, 321–350.

RAYSON, P. *et al.*, 2005: VARD versus Word: A comparison of the UCREL variant detector and modern spell checkers on English historical corpora, in: *Proceedings of the Corpus Linguistics 2005 conference*, Birmingham, UK, July 14–17 2005.

ROBERTSON, A. M. & P. WILLETT, 1993: A comparison of spelling-correction methods for the identification of word forms in historical text databases, in: *Literary and Linguistic Computing* 8, 143–152.

RUSSELL, R. C., 1918: Soundex coding system, in: *United States Patent* 1,261,167.

VITERBI, A. J., 1967: Error bounds for convolutional codes and an asymptotically optimal decoding algorithm, in: *IEEE Transactions on Information Theory* 13, 260–269.