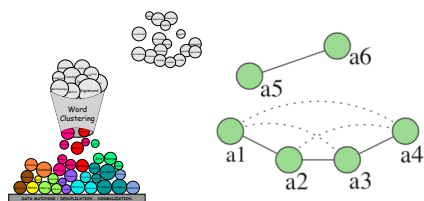


Poloautomatická normalizace slov z matričních záznamů

David Hříbek*



Abstrakt

V této práci je řešeno shlukování slov z matričních záznamů tak, aby v každém shluku byla slova, která jsou si navzájem velmi podobná.

V řešení bylo použito procesu deduplikace záznamů, který byl pozměněn tak, aby odpovídal požadavkům práce. Díky roztřídění záznamů do shluků bylo následně možné tyto shluky normalizovat, tedy přidělit celým shlukům nebo jednotlivým slovům jejich normalizované názvy.

V rámci práce byla přidána možnost normalizace slov do již existující webové aplikace, pro správu matričních záznamů, DEMoS.

Díky normalizaci slov jsou matriční záznamy lépe čitelné a je mezi nimi možné snadněji vyhledávat.

Klíčová slova: Data-Matching — Deduplication — Parish Records — Normalization

Přiložené materiály: DEMoS

*xhribe02@stud.fit.vutbr.cz, Faculty of Information Technology, Brno University of Technology

1. Úvod

V minulosti byly matriční záznamy zapisovány pouze do matričních knih (matrik). V dnešní době se často matriční záznamy digitalizují a ukládají do elektronických databází, díky čemuž je možné mezi těmito záznamy snadněji vyhledávat a třídit je.

Matriční záznamy se zapisují do knih, které mají pevně danou strukturu (tabulka s pojmenovanými sloupci). Každý záznam mohl být ale zapsán jiným člověkem, což má za následek, že některé záznamy mohou být obsáhlejší, či méně čitelné, než jiné. Mohou také obsahovat chyby, překlepy nebo různé varianty stejného slova. Během staletí se taktéž mohla změnit podoba některých slov, která jsou do matrik zapisována. Například název města, který se kdysi zapisoval jako „Bozsko-wycze“ je dnes zapisován jako „Boskovice“. Totéž platí o křestních jménech, příjmeních, povoláních, atd. Podoba slov v matričních záznamech také záleží na

době zápisu. Během let se na stejném místě mohly zapisovat matriční knihy v různých jazycích jako např. čeština 17. a 18. století, latina, němčina nebo moderní čeština. Tímto vznikají různé variace slov téhož významu, což ztěžuje možnost vyhledávání mezi záznamy.

V této práci je řešeno rozšíření již existující webové aplikace DEMoS, která se zabývá správou matričních záznamů, o možnost normalizace jednotlivých slov v těchto záznamech. Normalizací slova je myšleno přiřazení danému slovu jeho dnešní podoby zápisu. Pro slova 'Josephus' a 'Joseph' by dnešní podoba zápisu byla například 'Josef'. Díky normalizaci je zvýšena efektivita vyhledávání a čitelnost matričních záznamů. Mezi typy slov, která jsou normalizována patří například jména, příjmení, povolání, obce a další.

Řešení tohoto problému se skládá z dvou částí – nalezení slov, která jsou si navzájem podobná (např.

19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36

jména 'Petr' a 'Peter' nebo 'Františka' a 'Francziska') a jejich roztrídění do shluků s následným přiřazením normalizovaného názvu jednotlivým slovům ve shlucích nebo celým shlukům. V první části je využit proces porovnávání dat (*Data-Matching*). Jelikož jsou hledány podobné záznamy v rámci jedné databáze, jedná se o speciální případ procesu porovnávání dat – deduplikace (*Deduplication*). Výsledkem tohoto procesu je, že každé slovo určené k normalizaci (jméno, příjmení, obec, povolání, atd.) je označeno číslem shluku, do kterého patří. Jednotlivým slovům je pak možné pomocí webové aplikace DEMoS přiřadit normalizovanou variantu a díky vytvořeným shlukům slov je možné tyto normalizované varianty sdílet v rámci daného shluku a nabízet tak doporučené normalizované varianty.

53 2. Použité technologie

54 Webové aplikace DEMoS je napsána v jazyce PHP
55 (Nette Framework). Algoritmus pro shlukování slov
56 je napsán v jazyce Python. Normalizace slov během
57 vytváření nebo úpravy záznamů je řešena asynchronně,
58 je zde využita technologie AJAX a serializační formát
59 JSON. Dále je použit programovací jazyk JavaScript a
60 knihovna JQuery.

61 3. Proces shlukování záznamů

62 Tato část bude věnována využití shlukové analýzy
63 při procesu shlukování záznamů. Bude zde čerpáno
64 z knihy [2] od autora Petera Christena, která vychází
65 z více než 300 článků na toto téma a shrnuje tak celou
66 problematiku. Pokud není řečeno jinak, je v této části
67 čerpáno z této knihy.

68 3.1 Předzpracování záznamů

69 Předzpracování záznamů [2] (*Data Pre-Processing*) je
70 důležitou součástí shlukové analýzy. Data, která vs-
71 tupují do procesu shlukování se mohou lišit ve formátu
72 a struktuře, a také mohou obsahovat chyby. Hodnoty
73 záznamů nemusí být vždy zadány v prvním pádu. Navíc
74 mohou obsahovat nadbytečná slova, která nepředstavují
75 informace užitečné pro shlukování. Mezi taková slova
76 patří například předložky. To, jakým způsobem byla
77 data pořízena také ovlivňuje jejich kvalitu. Data mohla
78 být přepisována z hůře čitelných zdrojů nebo naskenová-
79 na a získána pomocí optického rozpoznávání znaků.
80 Chyby při zaznamenávání dat také nastávají pokud
81 byla data diktována.

82 Z těchto důvodů je nutné před začátkem samotného
83 shlukování data pročistit a standardizovat. Pročištěním
84 je zvýšena efektivita shlukování. Proces čištění dat je

znám jako *Data Cleansing*. V následujícím seznamu
jsou uvedeny příklady čištění dat:

- **Kódování znaků.** Dekódování znaků v kódová-
ní UNICODE do kódování ASCII. Pro každý
znak v kódování UNICODE je nalezen nejvíce
podobný znak v kódování ASCII. Tento krok
je důležitý, pokud data mohou obsahovat znaky
z cizích jazyků.
- **Jednotná velikost písmen.** Informace o ve-
likosti písmen je možno při shlukování zanedbat.
Tato informace není při shlukování rozhodující.
Je proto vhodné změnit velikost písmen na velká,
nebo malá písmena.
- **Nadbytečné znaky.** Znaky, které nejsou důleži-
té pro shlukování (nepřináší žádnou novou in-
formaci) je nutné odstranit. Mezi takové znaky
patří například všechny nealfanumerické znaky,
kromě mezer. Mezery oddělují jednotlivá slova,
tuto informaci je nutné zachovat.
- **Sekvence opakujících se znaků.** Sekvence
opakujících se znaků také většinou nepřinášejí
informace, které by se měly zohledňovat při
porovnávání záznamů. Například slova 'auto' a
'auutoooo' pravděpodobně představují stejnou
entitu, auto. Proto je nutné nahradit sekvence
opakujících se znaků jedním znakem.
- **Nahrazení znaků.** Podobnost některých znaků
je větší než podobnost jiných znaků. Znaky
jako y a i nebo w a v jsou si velice podobné, a
proto je možné vybrat jeden z těchto znaků a
nahradit jím všechny výskyty druhého znaku.
V závislosti na typu dat je možné sestavit více
pravidel pro nahrazení znaků nebo sekvencí zna-
ků, a uložit je do vyhledávací tabulky (*lookup
table*)
- **Číslice.** Pokud shlukujeme například názvy obcí
nebo jména osob, číslice je také možno odstranit.
- **Nadbytečná slova.** Hodnoty záznamů mohou
obsahovat slova, která nepřinášejí žádnou infor-
maci, užitečnou pro shlukování – „stop slova“
(*stop words*¹) – mezi taková slova patří běžně
užívaná slova v různých jazycích.
Existují předpřipravené slovníky se „stop slovy“
pro různé jazyky. V českém jazyce mezi stop
slova patří například: *a, aby, byl, byla, pro, vsak,
ktere*.

Proces čištění dat je závislý na typu dat. Je proto
nutné zvolit způsob čištění dat tak, aby daná data
nebyla znehodnocena. Důležité také je, aby při pro-
cesu čištění dat nebyla ztracena původní vstupní data.

¹stop words - https://en.wikipedia.org/wiki/Stop_words

135 3.2 Indexace

136 Indexace [2] je důležitým krokem před samotným
137 porovnáváním záznamů. Pročištěná a standardizo-
138 vaná data jsou nyní připravena na porovnávání. Při
139 shlukování dvou databází by tedy bylo nutné porov-
140 nat každý záznam z první databáze s každým záznamem
141 z druhé databáze. Pokud je shlukována jedna databáze,
142 musí být každý záznam porovnán se všemi ostatními
143 záznamy v dané databázi. Proces porovnávání má tedy
144 kvadratickou časovou složitost. Při každém porovnání
145 dvojice záznamů je navíc nutné porovnat všechny atri-
146 buty obou záznamů. Porovnávání je nejnáročnější
147 operací procesu shlukování záznamů.

148 Když je porovnáván jeden záznam se všemi os-
149 tatními záznamy, většina porovnání dopadne neúspěšně.
150 Jen malé procento skončí úspěchem (záznamy předsta-
151 vují stejnou entitu). Aby bylo zredukováno potenciální
152 velké množství kandidátních párů k porovnání, je nutné
153 vybrat ty páry, u kterých je velká pravděpodobnost na
154 úspěch porovnání. Tyto odfiltrované páry jsou poté
155 předány k detailnímu porovnání. Proces vybírání kan-
156 didátních párů k detailnímu porovnání se nazývá in-
157 dexace.

158 Proces indexace má kvadratickou časovou složitost,
159 jelikož je nutné porovnat všechny záznamy podle dané-
160 ho atributu mezi sebou. Využití indexace tedy nemá
161 smysl, pokud shlukujeme záznamy na základě jed-
162 noho atributu. Indexace může naopak výrazně zmenšit
163 časovou složitost porovnávání záznamů, pokud záznamy
164 obsahují více atributů, které chceme zohlednit při po-
165 rovnávání.

166 3.3 Porovnání záznamů

167 Dalším krokem v procesu shlukování záznamů je samo-
168 tné porovnávání [2] jednotlivých záznamů mezi se-
169 bou (*Record Comparison*). Existují tři přístupy pro
170 porovnávání – přesné porovnání, částečné (*truncate*) a
171 porovnání s využitím kódovací funkce.

Pro porovnání jednotlivých atributů záznamů jsou
využity porovnávací funkce (*comparison functions*).
Obecný předpis porovnávací funkce zobrazuje rovnice
(1). Jedná se o funkci pro výpočet podobnosti.

$$s = \text{sim}(a_i, a_j) \quad (1)$$

172 Parametry a_i a a_j jsou atributy pro porovnání, výsled-
173 kem je hodnota s , která určuje jak moc jsou si tyto
174 atributy podobné. Pro podobnost s platí:

- 175 • Pokud $\text{sim}(a_i, a_j) = 1$, atributy jsou zcela shodné.
- 176 • Pokud $\text{sim}(a_i, a_j) = 0$, atributy jsou zcela odlišné.
- 177 • Pokud $0 \leq \text{sim}(a_i, a_j) \leq 1$, atributy jsou si do
178 jisté míry podobné.

To, co znamená, že jsou si atributy zcela shodné, zcela
odlišné nebo do jisté míry podobné záleží na konkrétní
použité funkci podobnosti. Lze také využít funkce
pro výpočet vzdálenosti mezi atributy. Podobnost a
vzdálenost jsou navzájem převoditelné a platí, že čím
jsou si atributy podobnější, tím je jejich vzdálenost
menší. Čím je vzdálenost atributů větší, tím méně jsou
si podobné.

Výsledkem přesného porovnávání je buď hodnota
1 (atributy jsou zcela shodné) nebo hodnota 0 (atributy
jsou zcela odlišné). Funkce podobnosti pro přesné
porovnávání je zobrazena v následující rovnici:

$$\text{sim}_{\text{exact}}(a_1, a_2) = \begin{cases} 1.0 & \text{pokud } a_1 = a_2 \\ 0 & \text{pokud } a_1 \neq a_2 \end{cases} \quad (2)$$

Variantou přesného porovnávání je částečné porov-
návání, kdy jsou přesně porovnávány pouze části obou
atributů (prefix nebo suffix). V tomto případě musí být
oba atributy typu řetězec.

Pro zjištění míry podobnosti dvou řetězců je nutné
využít funkce pro výpočet podobnosti řetězců. Některé
z těchto funkcí jsou popsány v následujícím seznamu.

- **Editační vzdálenost.** Nejznámější implemen-
tací editační vzdálenosti je Levenshteinova edita-
ční vzdálenost. Tato vzdálenost je definována
jako nejmenší možný počet operací *vložení znaku*,
odstranění znaku a *substituce znaku* pro změnu
jednoho řetězce na druhý řetězec. Pro výpočet
vzdálenosti je využito dynamické programování².
Časová složitost je kvadratická. Metoda ne-
zohledňuje ve které části jsou řetězce odlišné
– odlišnosti řetězců na začátku nebo na konci
přikládá stejnou váhu – hodí se proto pro krátké
řetězce nepředstavující jména nebo názvy. Ukáz-
ka výpočtu Levenshteinovy editační vzdálenosti
pro dva páry jmen je zobrazena v následujících
tabulkách.

Podobnost dvou vstupních řetězců s_1 a s_2 lze po-
mocí Levenshteinovy vzdálenosti vyjádřit jako:

$$\text{sim}_{\text{Levenshtein}}(s_1, s_2) = 1.0 - \frac{\text{dist}_{\text{Levenshtein}}(s_1, s_2)}{\max(|s_1|, |s_2|)},$$

kde $\text{dist}_{\text{Levenshtein}}$ je Levenshteinova vzdálenost,
 $|s_1|$ délka řetězce s_1 a $|s_2|$ délka řetězce s_2 .
Rozšířením této metody je editační vzdálenost
Demerau-Levenshtein, která přidává čtvrtou
operaci – výměnu dvou sousedních znaků. Vý-
měna dvou sousedních znaků je častou chybou
při získávání dat.

²Dynamické programování - https://cs.wikipedia.org/wiki/Dynamick%C3%A9_programov%C3%A1n%C3%AD

		0	1	2	3	4	5
			p	a	v	l	a
0		0	1	2	3	4	5
1	p	1	0	1	2	3	4
2	a	2	1	0	1	2	3
3	v	3	2	1	0	1	2
4	e	4	3	2	1	2	3
5	l	5	4	3	2	1	2

		0	1	2	3	4	5
			r	a	d	i	m
0		0	1	2	3	4	5
1	r	1	0	1	2	3	4
2	a	2	1	0	1	2	3
3	d	3	2	1	0	1	2
4	k	4	3	2	1	2	3
5	a	5	4	3	2	3	4

Tabulka 1. Ukázka výpočtu Levensthteinovy editační vzdálenosti dvou jmen. V první tabulce je vypočtena vzdálenost mezi jmény 'pavla' a 'pavel' rovna 2. V druhé tabulce je vypočte vzdálenost mezi jmény 'radim' a 'radka' rovna 4. Tučné číslice vyznačují cestu ke konečnému výsledku. Tyto tabulky byly převzaty a upraveny z knihy [2].

- **Q-gram.** Principem této metody, která je také nazývána N-gram, je rozdělit oba vstupní řetězce na krátké podřetězce o délce q (Q-gram řetězce). Délka těchto řetězců je obvykle 2 (*bigram*) nebo 3 (*trigram*) znaky. K rozdělení vstupních řetězců na Q-gram řetězce je použit princip posuvného okna o délce q , které je posunováno od začátku řetězce na jeho konec. V každém kroku je vytvořen jeden Q-gram řetězec. Počet vzniklých Q-gram řetězců odpovídá $c = |s| - q + 1$, kde $|s|$ je délka řetězce a q je délka Q-gram řetězců. Podobnost vstupních řetězců s_1 a s_2 je pak možné vypočítat pomocí některého z následujících koeficientů, kde c je počet společných Q-gram řetězců, které se nacházejí v obou vstupních řetězcích, a c_1 a c_2 je počet Q-gram řetězců vzniklých z řetězce s_1 resp. s_2 :

$$sim_{overlap}(s_1, s_2) = \frac{c}{\min(c_1, c_2)} \quad (3)$$

$$sim_{jaccard}(s_1, s_2) = \frac{c}{(c_1 + c_2) - c} \quad (4)$$

$$sim_{dice}(s_1, s_2) = \frac{2 \times c}{c_1 + c_2} \quad (5)$$

Možným rozšířením metody Q-gram je zohlednění pozice jednotlivých Q-gram řetězců v rámci vstupního řetězce. Při určení počtu společných Q-gram řetězců nejsou brány v úvahu Q-gram řetězce, které jsou od sebe vzdáleny více, než je maximální povolená vzdálenost. Výpočet podobnosti řetězců pomocí metody Q-gram je možné využít pro krátké řetězce, které nepředstavují jména. Následující tabulka zobrazuje rozdělení vstupního řetězce na jednotlivé Q-gram řetězce:

- **Jaro, Jaro-Winkler.** Tyto metody byly speciálně vyvinuty pro výpočet podobnosti jmen. Kombinují editační vzdálenost s podobnými postupy

Řetězec	Bigram	Trigram	Bigram s pozicí
david	da, av, vi, id	dav, avi, vid	(da,1), (av,2), (vi,3), (id,4)
pavel	pa, av, ve, el	pav, ave, vel	(pa,1), (av,2), (ve,3), (el,4)

Tabulka 2. Ukázka rozdělení vstupního řetězce na bigram řetězce, trigram řetězce a bigram řetězce s pozicí.

jako u podobnosti Q-gram. Výpočet podobnosti dvou řetězců pomocí základní metody **Jaro** je zobrazen v rovnici (6). Pro výpočet počtu shodných znaků c a počtu transpozic t (dvojice sousedních znaků, které se liší pořadím znaků, např. 'da' a 'ad'.) je využito posuvné okno o velikosti poloviny délky většího z řetězců. Toto okno je posunováno současně po obou řetězcích. V každém kroku posunutí je zaznamenán počet znaků, které jsou shodné v obou oknech a počet transpozic.

$$sim_{jaro}(s_1, s_2) = \frac{1}{3} \left(\frac{c}{|s_1|} + \frac{c}{|s_2|} + \frac{c-t}{c} \right) \quad (6)$$

Na základě zjištění studií, že větší počet chyb ve jménech se nachází uprostřed nebo na konci řetězce, bylo vytvořeno rozšíření **Jaro-Winkler**, které klade větší váhu odlišnostem řetězců na jejich začátku, než uprostřed nebo na konci.

- **Longest Common Substring.** Hlavní myšlenkou tohoto algoritmu je najít a odstranit nejdelší možné podřetězce, které jsou společné v obou vstupních řetězcích. Tento postup se opakuje tak dlouho, dokud délka nejdelších společných podřetězců je větší nebo rovna předem stanovené délce l_{min} (minimální délka je většinou 2 nebo 3 znaky). Suma délek (l_c) všech nalezených společných podřetězců je pak použita pro výpočet podobnosti mezi vstupními řetězci s_1 a s_2 . Lze využít jeden z jichž dříve zmíněných koeficientů:

$$sim_{lcs_overlap}(s_1, s_2) = \frac{l_c}{\min(c_1, c_2)}, \quad (7)$$

$$sim_{lcs_jaccard}(s_1, s_2) = \frac{l_c}{(|s_1| + |s_2|) - |l_c|}, \quad (8)$$

$$sim_{lcs_dice}(s_1, s_2) = \frac{2 \times l_c}{|s_1| + |s_2|}. \quad (9)$$

Časová složitost výpočtu sumy (l_c) všech délek nejdelších společných podřetězců je kvadratická. Nevýhodou tohoto algoritmu je, že vypočtená podobnost vstupních řetězců nemusí být vždy symetrická – pokud jsou vstupní řetězce prohozeny, výsledná míra podobnosti se může lišit.

Toto lze ukázat na výpočtu podobnosti mezi řetězci $s_1 = 'janj'$ a $s_2 = 'jnja'$ a $l_{min} = 2$. V prvním kroku bude jako nejdelší nalezený společný podřetězec nalezen řetězec 'ja', po jehož odstranění z obou řetězců vzniknou nové řetězce $s'_1 = 'nj'$ a $s'_2 = 'jn'$, dále už nejsou nalezeny žádné další společné podřetězce, tudíž $l_c = 2$ a $sim_{lcs} = 0.5$. Pokud jsou vstupní řetězce vyměněny, tedy $s_1 = 'jnja'$ a $s_2 = 'janj'$, jako první nejdelší podřetězec je nalezen řetězec 'nj', po jehož odstranění vzniknou nové řetězce $s'_1 = 'ja'$ a $s'_2 = 'ja'$. V dalším kroku je nalezen společný podřetězec 'ja', po jehož odstranění z obou řetězců už nejsou nalezeny další společné podřetězce, delší než 2 ($l_{min} = 2$), proto $l_c = 4$ a $sim_{lcs} = 1$. Pro vyřešení tohoto problému je nutné vypočítat podobnost vstupních řetězců v obou pořadích a výsledek zprůměrovat.

3.3.1 Fonetické kódování atributů

Cílem funkcí pro fonetické kódování řetězců [2] je převést vstupní řetězec na kód, který odpovídá tomu, jak by daný řetězec byl vysloven. Tyto funkce jsou závislé na konkrétním jazyce, existují proto varianty těchto funkcí pro konkrétní jazyky. Fonetické kódování je možné použít také při výpočtu podobnosti mezi řetězci, kdy jsou nejdříve jednotlivé řetězce foneticky zakódovány a poté jsou tyto kódy porovnány. Fonetické kódování se používá u jmen, příjmení a názvů. Nejznámější algoritmus pro fonetické kódování se nazývá Soundex [3]. Od tohoto algoritmu jsou poté odvozeny další varianty jako např. Phonex, Phonix nebo Double-Metaphone.

Soundex Výsledný kód vznikne převodem vstupního řetězce (např. jméno) na kód ve tvaru počátečního písmene řetězce, následovaného několikamístným číselným kódem (většinou je použit 3 místný číselný kód). Tento číselný kód je vytvořen pomocí transformační tabulky. Tabulka popisuje transformace znaků na číslice. Z výsledného kódu jsou odebrány sekvence duplicitních číslic a je zarovnán na požadovanou délku. Pokud je délka kódu menší než požadovaná délka kódu, kód je doplněn zprava nulami. Pokud je délka kódu větší než požadovaná délka kódu, výsledný kód je ořezán na požadovanou délku.

Double-Metaphone Hlavní nevýhodou předešlého algoritmu je jeho závislost na konkrétním jazyce. V reálných databázích se nenacházejí jen anglické názvy (např. jména), proto se algoritmus Double-Metaphone zaměřuje nejen na anglická, ale také evropská a asijská jména a názvy. Vhodný je také pro data, která obsahují chyby.

Prvním krokem je předzpracování vstupního řetězce, kde je aplikováno velké množství pravidel pro modifikaci řetězce. Rozdílem tohoto algoritmu oproti předešlému je, že výsledný kód neobsahuje číslice, pouze znaky. Pro některá jména jsou navíc navraceny dva výsledné fonetické kódy, místo jednoho. Tyto dva kódy vzniknou odlišným pořadím aplikování fonetických transformačních pravidel.

3.4 Klasifikace

Záznamy, které byly vyfiltrovány v kroku Indexace a detailně porovnány v předešlém kroku Porovnání záznamů je nyní možné klasifikovat [2]. Záznamy lze klasifikovat do 3 tříd: shodné (*Matches*), neshodné (*Non-matches*) a potenciálně shodné (*Potential matches*).

Výsledkem předešlého kroku Porovnání záznamů je pro každou porovnanou dvojici záznamů vektor, který obsahuje míru podobnosti jednotlivých atributů záznamů. Ukázka výstupu porovnání dvou dvojic záznamů, které obsahují atributy jméno a příjmení, je zobrazena v následující tabulce.

Záznam	Jméno	Příjmení	Suma podobnosti
a1	Petr	Novák	
b1	Peter	Nový	
	0.88	0.66	1.54
a2	Jan	Šimeček	
b2	Johan	Schimetzek	
	0.75	0.58	1.33

Tabulka 3. Ukázka porovnání dvou dvojic záznamů. Tučná čísla představují vzniklé vektory, které obsahují míru podobnosti jednotlivých atributů.

Tento vektor je dále použit při klasifikaci. Základním způsobem klasifikace záznamů do tříd je **Prahování**. Využívají se zde dva prahy t_{up} a t_{low} , a suma podobnosti záznamů (*SimSum*). Klasifikace do tříd pak probíhá podle těchto pravidel:

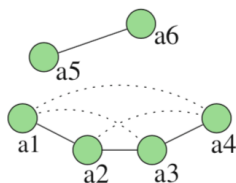
$$\begin{aligned} SimSum[r_i, r_j] &\geq t_{up} \implies [r_i, r_j] \rightarrow Match, \\ t_{low} \leq SimSum[r_i, r_j] &\leq t_{up} \implies [r_i, r_j] \rightarrow Potential Match, \\ SimSum[r_i, r_j] &\leq t_{low} \implies [r_i, r_j] \rightarrow Non - Match, \end{aligned}$$

kde $SimSum[r_i, r_j]$ je suma podobnosti záznamů r_i a r_j . Pokud je hodnota prahů stejná, tedy $t_{up} = t_{low} = t$, třída Potenciální shoda je odstraněna. V takovém případě se jedná se o klasifikaci do dvou tříd. Klasifikace do dvou tříd pak probíhá podle těchto pravidel:

$$\begin{aligned} SimSum[r_i, r_j] &\geq t \implies [r_i, r_j] \rightarrow Match, \\ SimSum[r_i, r_j] &< t \implies [r_i, r_j] \rightarrow Non - Match. \end{aligned}$$

363 3.5 Transitivní uzavření

364 Každý záznam může být klasifikován jako shodný
365 s více jinými záznamy. V takovém případě vzniká
366 řetězec záznamů, které jsou si podobné. Pokud nas-
367 tane situace, že dvojice záznamů (a_2, a_1) a (a_2, a_3)
368 byly klasifikovány jako shodné (Matches) a dvojice
369 záznamů (a_1, a_3) byla klasifikována jako neshodná
370 (Non-Matches), je nutné řešit transitivní uzavření těchto
371 záznamů. Příklad takové situace je zobrazen na obrázku 1.



Obrázek 1. Ukázka klasifikace záznamů do dvou tříd (Matches, Non-Matches). Jednotlivé záznamy jsou zobrazeny jako zelené kruhy. Plnou čarou jsou propojeny záznamy, které byly klasifikovány jako shodné (Matches), čárkovanou čarou jsou propojeny záznamy, které nebyly klasifikovány jako shodné (Non-Matches). Obrázek převzat z knihy [2].

372 Pokud však záznam a_2 představuje stejnou entitu
373 (objekt reálného světa) jako záznamy a_1 a a_3 , tak pak
374 i záznam a_1 musí představovat stejnou entitu jako
375 záznam a_3 . Řešením takové situace je klasifikace dvo-
376 jice záznamů (a_1, a_3) jako shodné (Matches).
377

378 4. Testování kvality shlukování

379 Posledním krokem procesu shlukování záznamů je
380 měření kvality shlukování [2]. K měření kvality shlu-
381 kování je nutné znát u každé porovnávané dvojice
382 záznamů referenční výsledek klasifikace (*true match*
383 *status*), který určuje, zda oba záznamy reálně předsta-
384 vují stejnou entitu (Matches) nebo ne (Non-Matches).
385 Takový datový set je pak nazýván *ground-truth*.

386 Každá dvojice záznamů pak lze zařadit do jedné
387 z následujících kategorií:

- 388 • **True positives (TP).** Záznamy, které byly klasi-
389 fikovány jako shodné a jsou opravdu shodné.
- 390 • **False positives (FP).** Záznamy, které byly klasi-
391 fikovány jako shodné, ale nejsou shodné.
- 392 • **True negatives (TN).** Záznamy, které byly klasi-
393 fikovány jako neshodné a jsou opravdu neshodné.
- 394 • **False negatives (FN).** Záznamy, které byly klasi-
395 fikovány jako neshodné, ale jsou shodné.

396 Cílem shlukování záznamů je, aby počty dvojic
397 v kategoriích TP a TN byly co nejvyšší, a zároveň
398 počty dvojic záznamů v kategoriích FP a FN byly

co nejnižší. Na základě počtů dvojic záznamů v jed- 399
notlivých kategoriích lze vypočítat různé metriky pro 400
měření kvality shlukování. Nejpočetnější kategorií je 401
kategorie TN, která většinou několikanásobně převyšuje 402
ostatní kategorie. Z tohoto důvodu nejsou pro určení 403
kvality shlukování vhodné metriky, které zohledňují 404
počty záznamů v této kategorii. 405

V následujícím seznamu jsou posány 3 metriky pro 406
měření kvality shlukování. Výsledkem všech těchto 407
metrik je hodnota v intervalu 0 až 1. 408

- **Precision.** Precision neboli přesnost určuje jaký 409
podíl z dvojic klasifikovaných jako shodné (TP, 410
FP) byl správně klasifikován jako shodný (TP). 411
Tento vztah zobrazuje následující rovnice: 412

$$precision = \frac{TP}{TP + FP}. \quad (10)$$

- **Recall.** Recall určuje jaký podíl z dvojic, které 413
jsou opravdu shodné (TP, FN) byl klasifikován 414
jako shodný (TP). Tento vztah je zobrazuje ná- 415
sledující rovnice: 416

$$recall = \frac{TP}{TP + FN}. \quad (11)$$

- **F-measure.** Výsledkem této metriky je harmon- 417
ický průměr předešlých dvou metrik. Hodnota 418
této metriky je vysoká pouze pokud je vysoká 419
hodnota metriky Precision i Recall. Vztah pro 420
výpočet této metriky zobrazuje následující rov- 421
nice: 422

$$f - measure = 2 \times \frac{precision \times recall}{precision + recall}. \quad (12)$$

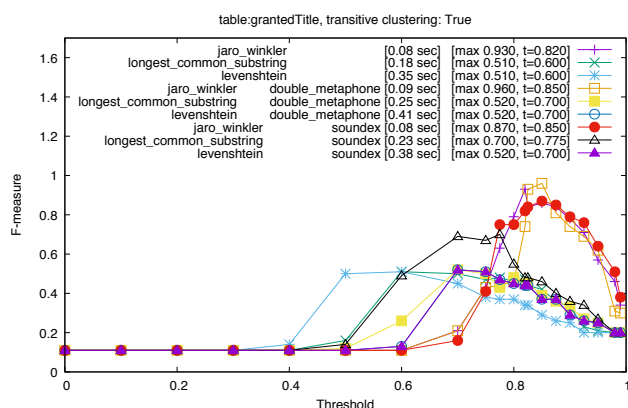
Tato metrika je často využívána pro měření kval- 423
ity shlukování. 424

4.1 Výběr algoritmů na základě testování 425

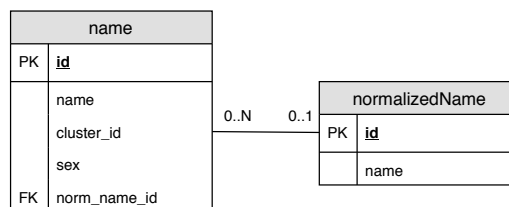
Parametry a výsledky testování shlukování jsou zo- 426
brazeny v tabulce 4. Pro testování byly použity data 427
z reálných maticových záznamů, které jsou psány češti- 428
nou 17. a 18. století, latinsky, německy a moderní 429
češtinou. Úspěšnost shlukování je posouzena vzhle- 430
dem k referenčnímu rozřazení slov do shluků (*ground-* 431
truth) pomocí metriky *F-measure*, která je popsána 432
níže. Finální použité funkce podobnosti řetězců, funkce 433
pro fonetické kódování atributů a práh podobnosti 434
byly vybrány na základě testování různých kombinací 435
těchto parametrů pro každý typy slov. Ukázka výsledků 436
testování pro tituly farářů je zobrazena v grafu na 437
obrázku 2. 438

Typ slova	Fon. Kód.	Funkce podobnosti řetězců	Práh	Počet slov	Úspěšnost shlukování (F-measure)
Jméno		JW	0.90	1573	~ 0.87
Příjmení	S	LCS	0.85	2470	~ 0.84
Povolání	DM	LCS	0.85	654	~ 0.96
Obec		JW	0.90	995	~ 0.92
Titul faráře	DM	JW	0.85	85	~ 0.96

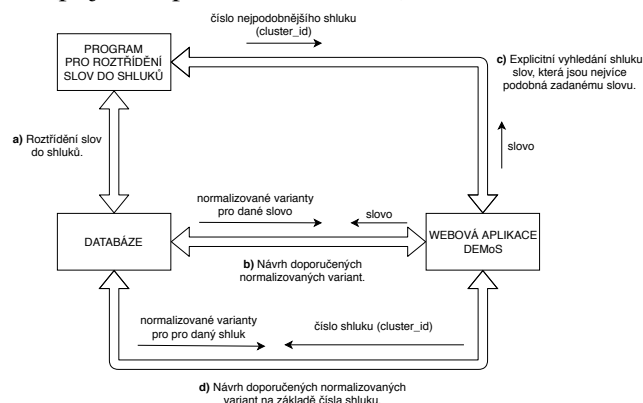
Tabulka 4. Parametry nejlepších naměřených výsledků tesování pro jednotlivé typy slov. Význam zkratk: S (Soundex), DM (Double-Metaphone), JW (Jaro-Winkler), LCS (Longest Common Substring).



Obrázek 2. Graf závislosti kvality shlukování na práhu podobnosti pro slova typu titul faráře. Pro každou kombinaci zvolené funkce podobnosti řetězců a fonetického kódování je vykreslena samostatná křivka. Při kombinaci fonetického kódování a funkce podobnosti řetězců je výsledná podobnost vypočtena poměrem 65:35. Z toho grafu lze zjistit, že nejlepší kombinací parametrů shlukování pro slova tituly faráře je kombinace Jaro-Winkler, Double Metaphone a práh podobnosti 0.85.



Obrázek 3. Schéma databáze pro uložení jmen z matričních záznamů a jejich normalizovaných variant. Toto schéma je aplikováno i pro ostatní typy dat (příjmení, povolání, obce, atd.).



Obrázek 4. Schéma obecného principu normalizace slov z matričních záznamů.

Originální slova z matričních záznamů budou nej- 455
dříve roztržena do shluků (část a na obrázku 4), po- 456
dle své podobnosti s ostatními slovy. K tomuto bude 457
využita shluková analýza a proces shlukování záznamů, 458
popsaný v kapitole č. 3. Výsledkem roztržení záznamů 459
do shluků bude u každé entity, z entitní množiny s ori- 460
ginálními slovy, vyplněné číslo shluku v atributu clus- 461
ter_id. Entity se stejným číslem shluku budou poté 462
chápány jako shodné (představující stejný objekt re- 463
álného světa — různé varianty stejného jména, po- 464
volání, příjmení, atd.). Proces roztržení slov do shluků 465
bude spuštěn automaticky v určitém časovém inter- 466
vale (např. každý den v nočních hodinách) pomocí 467
softwarového démona CRON³. 468

Při vytváření nebo editaci záznamů, pomocí we- 469
bové aplikace DEMoS, budou uživateli po zadání kon- 470
krétního slova nabídnuty (v období formulářového 471
výběrového prvku) doporučené normalizované vari- 472
anty pro daný typ a pohlaví slova. Tyto varianty budou 473
získány porovnáním zadaného slova s již uloženými 474
originálními slovy v databázi. Pokud zadané slovo již 475
existuje v databázi (část b na obrázku 4), budou prior- 476
itně doporučeny normalizované varianty, které již byly 477
použity u tohoto slova, následně pak normalizované 478
varianty slov, která se nacházejí ve shluku se stejným 479
číslem jako toto slovo. Pokud zadané slovo ještě neex- 480
istuje v databázi, uživatel bude mít možnost explicitně, 481

³CRON - <https://cs.wikipedia.org/wiki/Cron>

5. Návrh řešení normalizace

Shluky, vzniklé při procesu shlukování, jsou využity 440
pro normalizaci slov z matričních záznamů. Pro každý 441
typ slov, které je potřeba normalizovat, je vytvořena 442
databázová tabulka s původními slovy a tabulka s nor- 443
malizovanými variantami pro daný typ slov. Mezi 444
těmito dvěma tabulkami je vazba 1:N (jedno normal- 445
izované slovo může být přiřazeno k více originálním 446
slovům). Schéma databáze pro normalizaci jmen je 447
zobrazeno na obrázku 3. 448

5.1 Obecný princip normalizace slov

Princip navržené normalizace slov z matričních zá- 450
znamů je zobrazen na obrázku 4. Navržené řešení se 451
skládá ze čtyř částí, které jsou v obrázku označeny jako 452
část a, b, c a d. Následující text popisuje jednotlivé 453
části navrženého řešení. 454

482 pomocí tlačítka, požádat o vyhledání shluku slov (část
483 c na obrázku 4), které jsou nejvíce podobné s tímto
484 slovem. Z tohoto shluku budou poté nabídnuty normal-
485 izované varianty slov (část d na obrázku 4). Uživatel
486 bude mít vždy možnost zadat vlastní normalizovanou
487 variantu daného slova ručně. Výsledkem normalizace
488 slova je vyplnění atributu `norm_name_id` dané entity.

489 V případě, že dvěma slovům v odlišných shlucích
490 (odlišná hodnota atributu `cluster_id`) bude přidě-
491 lena stejná normalizovaná varianta (stejná hodnota
492 atributu `norm_name_id`), při následujícím plánova-
493 ném procesu rozřídění slov do shluků budou tyto
494 shluky spojeny v jeden shluk. Tímto je možné spojovat
495 vzniklé shluky a sdílet tak normalizované varianty.

496 5.2 Čas normalizace

497 Slova z matričních záznamů je možné normalizovat
498 již při vytváření nebo úpravě záznamů (ukázka zo-
499 brazena na obrázku 7). Pro hromadnou normalizaci
500 slov je vytvořena samostatná stránka, kde lze slova
501 filtrovat podle typu (jméno, příjmení, atd.), obsahu,
502 autra, pohlaví (pouze jména a příjmení) a stavu nor-
503 malizace (normalizováno/nenormalizováno). Ukázka
504 této stránky je zobrazena na obrázku 9.

505 6. Implementace

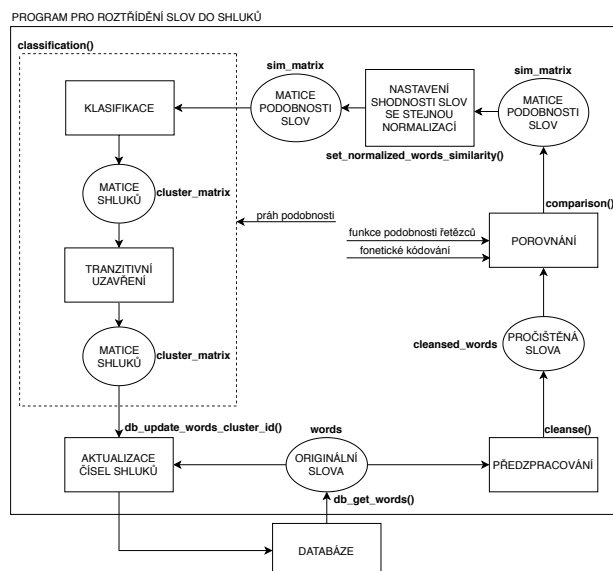
506 V této kapitole bude popsána implementace normal-
507 izace slov z matričních záznamů podle návrhu, který
508 byl popsán v předešlé kapitole. Implementace pro-
509 gramu pro rozřídění slov do shluků je částečně in-
510 spirována postupem [1].

511 Rozřídění slov do shluků

512 V této sekci je popsána část a) návrhu řešení, který je
513 zobrazen na obrázku 4. V rámci této části je implemen-
514 tován *Proces shlukování záznamů*, který byl popsán
515 v kapitole č. 3. Na obrázku 5 je zobrazeno schéma
516 programu pro rozřídění slov do shluků.

517 Pro rozřídění slov do shluků byly využity pouze
518 kroky *Předzpracování záznamů*, *Porovnání záznamů*
519 a *Klasifikace*. Jelikož je v této práci řešeno rozřídění
520 slov do shluků (= rozřídění záznamů do shluků na
521 základě jednoho atributu – slova), je nutné porovnat
522 každé slovo se všemi ostatními slovy. Z tohoto důvodu
523 nebyl využit krok *Indexace*.

524 Proces rozřídění slov do shluků, který je zobrazen
525 na obrázku 5 a byl implementován v rámci třídy *Words-*
526 *Clustering*, je spuštěn pomocí metody *execute*.
527 Nejdříve jsou z databáze načteny slova ke shlukování
528 (na základě předaných parametrů při instanciaci), která
529 jsou uložena v podobně seznamu do instanční proměnné
530 *words*. Podle počtu načtených slov jsou pomocí



Obrázek 5. Schéma programu pro rozřídění slov do shluků.

knihovny NumPy⁴, která umožňuje práci s maticemi, 531
vytvořeny čtvercové matice *sim_matrix* a *cluster-* 532
matrix. Tyto matice mají počet řádků i sloupců 533
daný počtem načtených slov a jsou inicializovány nu- 534
lami. 535

Načtená slova jsou poté předána ke předzpracování, 536
jehož výsledkem je seznam pročištěných slov, uložený 537
v instanční proměnné *cleansed.words*. Pročištěná 538
slova jsou následně porovnána invokací metody *com-* 539
parison. V této metodě je pro každou dvojici slov 540
vypočtena podobnost pomocí zvoleného fonetického 541
kódování a funkce podobnosti řetězců. Výsledek je 542
uložen do příslušné buňky matice podobnosti *sim-* 543
matrix (příklad vypočtené matice podobnosti je zo- 544
brazen v tabulce 5, vlevo). Jelikož je matice podob- 545
nosti symetrická, stačí vypočítat pouze část nad nebo 546
pod diagonálou, včetně. 547

Po vypočtení podobnosti pro každou dvojici slov 548
je matice podobnosti upravena tak, aby slova která 549
již jsou normalizována a zároveň jejich normalizace 550
jsou shodné (atribut *norm_name_id*), měla podob- 551
nost 1. Toto zajišťuje metoda *set_normalized-* 552
words_similarity (příklad matice podobnosti 553
před a po explicitním nastavením podobnosti slov je 554
zobrazen v tabulce 5). 555

Dalším krokem je klasifikace porovnaných dvojic 556
slov do tříd. Zde byla využita metoda prahování s jed- 557
notným prahem, tedy klasifikace do dvou tříd (shodné, 558
neshodné), která byla popsána v kapitole č. 3.4. Klasi- 559
fikaci zajišťuje metoda *classification*. Po in- 560
vokaci této metody je v matici podobnosti *sim_mat-* 561
rix porovnána podobnost každé dvojice s prahem 562

⁴NumPy - <http://www.numpy.org/>

563 podobnosti `threshold`. Pokud je daná podobnost
564 větší nebo rovna práhu podobnosti, do ekvivalentní
565 buňky matice `cluster_matrix` je uloženo číslo
566 řádku matice, na kterém se daná buňka nachází (příklad
567 vypočtené matice shluků `cluster_matrix` je zo-
568 brazen v tabulce 6, vlevo). Posledním krokem klasi-
569 fikace je tranzitivní uzavření, které je aktivováno na
570 základě parametru `transitive_closure`, při in-
571 stanci. Algoritmus pro tranzitivní uzavření jsem
572 narvhl tak, že matice `cluster_matrix` je prochá-
573 zena po sloupcích, od konce. Pro každý sloupec je
574 zjištěna nejmenší hodnota `col_min`, která se v něm
575 nachází. Pokud se v daném sloupci nacházejí i buňky
576 s hodnotou větší než nalezená nejmenší hodnota, jsou
577 hodnoty těchto buňek, včetně buňek na stejném řádku,
578 změněny na hodnotu `col_min`. Procházení matice
579 je ukončeno pokud při předešlém průchodu nebyla
580 změněna žádná buňka. Příklad původní matice shluků
581 a matice shluků po tranzitivním uzavření je zobrazen
582 v tabulce 6.

583 Posledním krokem procesu rozřídění slov do shlu-
584 ků je aktualizace čísel shluků v databázi. Toto zajišťuje
585 metoda `db_update_words_cluster_id`. Nejdříve
586 jsou z vypočtené matice shluků `cluster_matrix`
587 získány čísla shluků jednotlivých slov. Číslo shluku
588 odpovídá nejmenšímu číslu ve sloupci daného slova.
589 Získání čísel shluků pro jednotlivá slova je zobrazeno
590 v tabulce 6. U každého slova v seznamu `words` je pak
591 v databázi aktualizováno číslo shluku `cluster_id`.

592 Výpočet nejpodobnějšího shluku pro zadané 593 slovo

594 V této sekci je posána část c) návrhu řešení, který je zo-
595 brazen na obrázku 4. Cílem je nalézt pro zadané slovo
596 číslo shluku, ve kterém se nachází nejvíce podobné
597 slovo se zadaným slovem. Tento algoritmus implemen-
598 tuje metoda `get_best_cluster_id_of_unknown-`
599 `_word`. Vstupním parametrem je pouze zadané slovo
600 `word`, pro které má být nalezeno číslo nejpodobnější
601 shluk. Databázová tabulka, fonetické kódování, funkce
602 podobnosti řetězců a práh podobnosti jsou zadány při
603 instanci třídy `WordsClustering`. Tato metoda
604 je volána z webové aplikace DEMoS asynchronně, po-
605 mocí technologie AJAX⁵. Výsledkem je buď to číslo
606 nalezeného shluku, nebo v případě neúspěchu hodnota
607 -1.

608 Po invokaci této metody jsou pomocí již zmíněné
609 metody `db_get_words` načtena slova z databáze. Po-
610 užítím parametru `clustered_only` je zajištěno na-
611 čtení pouze slov, která již byla shlukována (mají vy-
612 plněný atribut `cluster_id`). Každé z těchto slov je

poté porovnáváno se zadaným slovem `word` s využitím
parametrů porovnávání, které byly zvolených při in-
stanci (pro každý typ slov byly zvoleny parametry
porovnávání na základě testování uvedeném v kapitole
č. 4). Pro slovo s největší podobností ke vstupnímu
slovu `word` je pak pomocí jazyka SQL⁶ nalezeno
v databázi číslo shluku `cluster_id`, jehož hodnota
je vrácena jako výsledek.

Normalizace slov při vytváření/editaci matričních záznamů

Součástí webové aplikace DEMoS již je formulář pro
vytvoření/editaci matričního záznamu. Tento formulář
obsahuje sekce pro zápis údajů o křestiteli, otci, matce,
kmotrech a dalších. Část údajů v těchto sekcích je
určena k normalizaci. Ukázka tří sekcí tohoto for-
muláře je zobrazena na obrázku 6.

Tento formulář byl rozšířen o normalizační část,
která se nachází na jeho konci. Ve výchozím stavu
je tato část formuláře skryta – lze ji zobrazit pomocí
tlačítka „Zobrazit standardizaci“, umístěném na konci
formuláře. Pro každé pole původního formuláře, které
je určeno k normalizaci, jsou v normalizační části
přidány čtyři formulářové prvky, pomocí kterých je
možné slovo v dané poli normalizovat:

- neměnné textové pole (`input type="text"`) s kopii zadaného slova určeného k normalizaci,
- výběrové pole (`select`) pro výběr doporučených normalizovaných variant, které již v databázi existují,
- textové pole pro možnost zadání vlastní normalizované varianty,
- tlačítko pro možnost explicitního návrhu normalizovaných variant (výpočet nejpodobnějšího shluku pro zadané slovo a návrh norm. variant použitých u slov v tomto shluku), které je zobrazeno pouze pokud zadané slovo nebylo nalezeno v databázi.

Tyto prvky jsou rovněž rozděleny do sekcí, podle
původního formuláře. Jednotlivá pole v normalizační
části jsou viditelná pouze pokud je původní pole for-
muláře vyplněno. Pokud není v sekci vyplněno žádné
pole, tato sekce je v normalizační části skryta. Na
obrázku 7 je zobrazena ukázka normalizační části for-
muláře pro již vyplněná pole původního formuláře na
obrázku 6.

Normalizace probíhá asynchronně na pozadí. Vše
je interaktivní – formulář se skrývá, odkrývá a doplňuje
v reálném čase podle aktuálně zadaných slov v polích
původního formuláře. Vyplněná normalizační část

⁵AJAX - <https://cs.wikipedia.org/wiki/AJAX>

⁶SQL - <https://cs.wikipedia.org/wiki/SQL>

662 formuláře je odeslána na server spolu s původním
663 formulářem při uložení. V případě doporučené nor-
664 malizované varianty je odeslán číselný identifikátor
665 dané varianty na kterou má být slovo normalizováno.
666 U vlastní normalizované varianty, je odesláno celé
667 slovo, na jehož základě je pak buďto vytvořena nová
668 normalizovaná varianta nebo nalezena již existující
669 varianta v databázi. Na tuto variantu je pak zadané
670 slovo normalizováno (její identifikátor je přiřazen do
671 atributu `norm_name_id` příslušného slova). Zadáni
672 vlastní normalizované varianty má větší prioritu než
673 výběr z doporučených normalizovaných variant.

674 Veškeré chování spojené s normalizací slov při
675 vkládání/editaci záznamů je naprogramováno v souboru
676 `normalizationAddEditRecord.js` pomocí již
677 zmíněného jazyka JavaScript a knihovny JQuery. Po
678 načtení stránky s formulářem jsou nalezeny pomocí
679 identifikátorů jednotlivá pole původního formuláře,
680 která jsou určena k normalizaci. Na tyto je navázána
681 událost `onChange`, která při zadání slova vyvolá jeho
682 předání do normalizační části formuláře, zobrazení/skry-
683 tí příslušné normalizační části formuláře a aktualizaci
684 doporučených normalizovaných variant pro zadané
685 slovo.

686 **Aktualizace doporučených normalizovaných vari-**
687 **ant při zadání slova** Jak už bylo řečeno, při zadání
688 slov do polí původního formuláře dojde k aktivaci
689 události `onChange` a tedy i k návrhu doporučených
690 normalizovaných variant. Toto je provedeno odesláním
691 zadaného slova a názvu daného formulářového pole
692 (atribut `name`) pomocí již zmíněných API endpoint url
693 na server, asynchronně použitím technologie AJAX.
694 Na straně serveru je invokována metoda `getNormal-`
695 `izedVariantsOfWordAndWordsInSameClu-`
696 `ster`, které je předáno zadané slovo a název daného
697 formulářového pole. Na základě názvu pole jsou pak
698 na serveru díky třídě `TableNames` zjištěny dodatečné
699 informace o tomto poli a poté vyhledány doporučené
700 normalizované varianty, které jsou pak odeslány zpět
701 ve formátu JSON. Po přijetí odpovědi je pak aktu-
702 alizováno výběrové pole s doporučenými variantami
703 v normalizační části formuláře pomocí jazyka JavaScript.
704 Jedná se o implementaci části **b** návrhu řešení, který
705 je zobrazen na obrázku 4.

706 **Aktualizace doporučených normalizovaných vari-**
707 **ant při explicitním požádání** Po zadání slova do
708 původního formuláře nemusí být vždy pro toto slovo
709 nalezeny doporučené normalizované varianty. Toto
710 může nastat v případě, že zadané slovo v databázi ještě
711 neexistuje nebo pro toto slovo nejsou k dispozici žádné
712 normalizované varianty. V prvním z těchto případů

je zobrazeno tlačítko „Navrhnout varianty“, pomocí 713
kterého má uživatel možnost explicitně požádat o návrh 714
doporučených normalizovaných variant. Jedná se o im- 715
plementaci částí **c** a **d** návrhu řešení, který je zobrazen 716
na obrázku 4. 717

Po kliknutí na toto tlačítko je zadané slovo a název 718
formulářového pole odeslán asynchronně na server po- 719
mocí zmíněných API endpoint url. Na straně serveru 720
dojde k invokaci metody `getNormalizedVari-` 721
`antsOfUnknownWord`, které je předáno zadané slo- 722
vo a název formulářového pole. Na základě názvu 723
pole jsou pak na serveru díky třídě `TableNames` 724
zjištěny dodatečné informace o tomto formulářovém 725
poli. Zadané slovo je pak porovnáváno se všemi slovy 726
daného typu v databázi, která jsou roztržďěná do shluků. 727
U slova, které je nejvíce podobné zadanému slovu a 728
zároveň tato podobnost je větší než stanovený práh 729
podobnosti, je zjištěno číslo shluku. Pomocí metody 730
`getNormalizedVariantsOfWordsInSameCl-` 731
`usterAsWord` jsou pak získány normalizované vari- 732
anty slov v tomto shluku, které jsou následně odeslány 733
zpět jako odpověď ve formátu JSON. Po přijetí od- 734
povědi je pak rovněž aktualizováno výběrové pole 735
s doporučenými normalizovanými variantami v nor- 736
malizační části formuláře pomocí jazyka JavaScript. 737

Proces porovnání zadaného slova s ostatními slovy 738
je implementován v rámci třídy `WordsClustering` 739
v souboru `WordsClustering.py` pomocí metody 740
`get_best_cluster_id_of_unknown_word`. Je- 741
likoř se nejedná o pouhý databázový dotaz, ale je nutné 742
spustit externí skript, není tato akce provedena auto- 743
maticky, ale až po kliknutí na zmíněné tlačítko. 744

Výběr normalizované varianty Po přijetí odpovědi 745
s normalizovanými variantami pro zadané slovo jsou 746
tyto varianty vloženy do příslušného formulářového 747
výběrového prvku v normalizační části formuláře. Od- 748
pověď ve formátu JSON obsahuje kolekci variant (`word`) 749
použitých u stejných slov jako zadané slovo a kolekci 750
variant (`cluster`) použitých u slov ve stejném shluku 751
jako zadané slovo. Nejdříve jsou vloženy varianty 752
z kolekce `word` (mají větší prioritu). Následně pak 753
do samostatné sekce varianty z kolekce `cluster`. 754
Jako poslední je vložena možnost „Nestandardizovat“, 755
pokud by uživatel nechtěl dané slovo normalizovat. 756

V závorce u každé doporučené varianty je zobrazen 757
počet použití dané varianty pro zadané slovo (varianty 758
v kolekci `word`) resp. pro slova ve stejném shluku 759
jako je zadané slovo (varianty v kolekci `cluster`). 760

Navržené varianty jsou primárně seřazeny podle 761
pohlaví – nejdříve jsou nabídnuty varianty použité 762
u slov stejného pohlaví jako je dané formulářové pole, 763
poté další přípustné varianty (např. u formulářového 764

765 pole pro mužské jméno budou nejdříve nabídnuty muž-
766 ské varianty a poté bezpohlavní varianty). Sekundární
767 seřazení je podle počtu použití. První doporučená
768 normalizovaná varianta tedy odpovídá nejvhodnější
769 normalizované variantě pro dané slovo a je vybrána
770 automaticky. Uživatel tedy nemusí normalizační část
771 formuláře otevírat pokud nechce – budou automat-
772 icky vybrány nejvhodnější normalizované varianty pro
773 zadaná slova.

774 Hromadná normalizace slov z matričních zá- 775 znamů

776 Hromadnou normalizaci slov je myšlena možnost vyp-
777 sání slov k normalizaci ze všech záznamů na jedné
778 stránce, kde tyto slova bude možné normalizovat. Sou-
779 částí této stránky je filtrační formulář, díky kterému
780 je možné vypsát pouze požadovaná slova. Slova lze
781 filtrovat podle typu (jména, příjmení, povolání, ...),
782 pohlaví, autora normalizace, stavu (normalizováno/ne-
783 normalizováno) a obashu. Aby nebylo vypsáno najed-
784 nou příliš mnoho slov, jsou slova rozříděna do stránek
785 mezi kterými lze přepínat. Počet slov na jedné stránce
786 lze také nastavit. Na obrázku 9 je zobrazena ukázka
787 stránky pro hromadnou normalizaci slov.

788 **Příprava slov pro hromadnou normalizaci** Proces
789 vykreslení této stránky je implementován metodou
790 `renderDefault`, která se nachází v presenteru `St-`
791 `andardizationPresenter`. Jelikož je pro ode-
792 slání filtračního formuláře zvolena metoda `GET`, pro-
793 tocolu `HTTP`⁷, parametry filtrování jsou předány v rámci
794 `URL`⁸ adresy. Tyto parametry jsou následně předány již
795 zmíněné metodě `getWordsForNormalization`,
796 která vrátí požadovaná slova pro hromadnou normal-
797 izaci. Pro každé z těchto slov jsou dodatečně získány
798 doporučené normalizované varianty voláním metody
799 `getNormalizedVariantsOfWordAndWordsIn-`
800 `nSameCluster`, které je předáno dané slovo, pohlaví
801 a název databázové tabulky, ve které je toto slovo
802 uloženo. Pro slova, která zatím nebyla normalizována,
803 je navíc pomocí metody `getMostFrequentWord-`
804 `OfCluster` zjištěno slovo které se nachází ve stejném
805 shluku jako zadané slovo a je nejčastěji používáné.
806 Toto slovo bude poté vypsáno ve sloupci `Standardi-`
807 `zováno` u slov, která zatím nebyla normalizována (au-
808 torem bude uveden počítač).

809 **Vykreslení slov pro hromadnou normalizaci** Jak
810 už bylo řečeno, slova pro hromadnou normalizaci jsou
811 vypsána pomocí tabulky, kde každý řádek představuje

jedno slovo určené k normalizaci. Po vykreslení stránky 812
již nejsou dodatečně doporučovány další normalizo- 813
vané varianty pomocí asynchronní komunikace se ser- 814
verem. 815

Interaktivní změny formuláře, které zahrnují zapín- 816
ání/vypínání tlačítka `Potvrdit` a výběrového prvku s do- 817
poručenými normalizovanými variantami, jsou im- 818
plementovány pomocí jazyka `JavaScript` v souboru 819
`standardizationPage.js`. V rámci toho soub- 820
oru je také implementována reakce na stisknutí tlačítka 821
`Potvrdit` – vybraná normalizovaná varianta resp. zadaná 822
vlastní normalizovaná varianta je odeslána asynchronně 823
na server pomocí zmíněných `API endpoint URL`, kde 824
je následně invokována metoda `updateNormaliz-` 825
`edVariantOfWord`. Této metodě je předán: iden- 826
tifikátor slova, které má být normalizováno, název 827
databázové tabulky ve které je toto slovo uloženo a 828
zadané normalizované varianty. Pokud uživatel zadal 829
vlastní normalizovanou variantu, je nalezena stejná 830
již existující normalizovaná varianta popř. vytvořena 831
nová normalizovaná varianta. Pokud uživatel nezadal 832
vlastní normalizovanou variantu, je k dispozici iden- 833
tifikátor vybrané normalizované varianty. Takto získaný 834
identifikátor normalizované varianty je pak uložen do 835
atributu `norm_name_id` záznamu tohoto slova, čímž 836
je slovo normalizováno. 837

7. Závěr 838

V této práci bylo řešeno rozšíření webové aplikace 839
`DEMoS`, pro správu matričních záznamů, o možnost 840
přřazení slovům v těchto záznamech jejich normal- 841
izovanou podobu a zvýšit tak čitelnost a efektivitu 842
vyhledávání těchto záznamů. Problém byl vyřešen 843
rozříděním slov do shluků, které následně slouží pro 844
sdílení normalizovaných variant slov v rámci daného 845
shluku. Při vytváření nebo editaci záznamů jsou pak 846
nabízeny uživateli normalizované varianty podle toho 847
v jakém shluku se slovo nachází, nebo by se nacházelo. 848

Shlukování slov bylo otestováno na reálných dat- 849
ech z matričních záznamů. Úspěšnost shlukování byla 850
posuzována vzhledem k referenčnímu rozložení slov 851
do shluků (*ground-truth*) a pohybovala se okolo 90 %. 852

Dalším krokem této práce by mohlo být například 853
přidání veřejně dostupné databáze povolání (`HISCO`) 854
a zvýšení úspěšnosti shlukování. 855

8. Poděkování 856

Rád bych poděkoval vedoucímu práce Ing. Jaroslavu 857
`Rozmanovi Ph.D.` za odbornou pomoc a informace 858
ohledně matričních záznamů, které mi během tvorby 859
této práce poskytl. 860

⁷`HTTP` - https://cs.wikipedia.org/wiki/Hypertext_Transfer_Protocol

⁸`URL` - https://cs.wikipedia.org/wiki/Uniform_Resource_Locator

- 862 [1] *Words clustering using Python Clustering list*
863 *of words*. Únor 2017, [Online; navštíveno
864 7.11.2018].
865 URL [https://rapply.weebly.](https://rapply.weebly.com/word-clustering/clustering-a-list-of-words/)
866 [com/word-clustering/](https://rapply.weebly.com/word-clustering/clustering-a-list-of-words/)
867 [clustering-a-list-of-words/](https://rapply.weebly.com/word-clustering/clustering-a-list-of-words/)
- 868 [2] Christen, P.: *Data Matching - Concepts and Tech-*
869 *niques for Record Linkage, Entity Resolution, and*
870 *Duplicate Detection*. Springer, 2012, ISBN 978-3-
871 642-31163-5.
- 872 [3] Zobel J., D. P.: *Phonetic string matching: Lessons*
873 *from information retrieval*. ACM SIGIR, Zürich,
874 Switzerland, 1996.

	Joan	Johanes	Johan	Petr	Peter		Joan	Johanes	Johan	Petr	Peter
Joan	1	0.72	0.88	0	0	Joan	1	1	0.88	0	0
Johanes	0	1	0.83	0.18	0.16	Johanes	0	1	0.83	0.18	0.16
Johan	0	0	1	0	0	Johan	0	0	1	0	0
Petr	0	0	0	1	0.88	Petr	0	0	0	1	0.88
Peter	0	0	0	0	1	Peter	0	0	0	0	1

Tabulka 5. Ukázka vypočtené matice podobnosti pro pět jmen. V levé tabulce je zobrazena původní vypočtená matice podobnosti. V pravé tabulce je zobrazena upravená matice podobnosti tak, aby slova která jsou již normlizována a jejich normalizace jsou shodné (v tomto případě se pro ukázkou jedná o jména Joan a Johanes), měla podobnost 1. Tučná čísla v pravé matici vznikla upravením původní matice podobnosti.

	Joan	Johanes	Johan	Petr	Peter		Joan	Johanes	Johan	Petr	Peter
Joan	1	0	1	0	0	Joan	1	0	1	0	0
Johanes	0	2	2	0	0	Johanes	0	1	1	0	0
Johan	0	0	3	0	0	Johan	0	0	1	0	0
Petr	0	0	0	4	4	František	0	0	0	4	4
Peter	0	0	0	0	5	Franta	0	0	0	0	4
číslo shluku	1	2	1	4	4	číslo shluku	1	1	1	4	4

Tabulka 6. Tabulka vlevo zobrazuje vypočtenou matici shluků (`cluster_matrix`) vzniklou z původní matice podobnosti, která je zobrazena v tabulce 5, vlevo. Tabulka vpravo zobrazuje původní matici shluků (vlevo) po tranzitivním uzavření. Díky tomu, že dvojice jmen (Joan, Johan) a (Johan, Johanes) byly klasifikovány jako shodné, po tranzitivním uzavření je i dvojice (Joan, Johanes) klasifikována jako shodná. Tučně jsou zvýrazněna čísla, která byla změněna při procesu tranzitivního uzavření. Poslední řádek obou tabulek obsahuje výsledná čísla shluků jednotlivých slov. Ty jsou získány jako nejmenější číslo ve sloupci.

^ Porodní bába

Jméno ?
Dorotha

Příjmení

Obec
Brun

Ulice

Číslo popisné

^ Dítě

Jméno ?

Příjmení
Holaubek

Vícerčata

Pohlaví
Muž

Lože
Manželské

Vyznání
katolík

Datum sňatku rodičů
dd/mm/yyyy

☐ Mrtvě rozené
☐ Nalezenec

^ Otec

☐ Mrtev ?

Jméno
Tobyass

Příjmení
Hříbek

Povolání
heytmán

Obec
Znoymenses

Ulice

Číslo popisné

Vyznání
-

Datum narození
dd/mm/yyyy

Obrázek 6. Ukázka části již existujícího formuláře pro vytvoření/editaci záznamu ve webové aplikaci DEMoS. Formulář je rozdělen do sekcí.

Standardizace

^ Porodní bába

Jméno	Doporučené	Vlastní
Dorotha	Dorota (3)	Vlastní varianta

Obec	Doporučené	Vlastní
Brun	Brno (2)	Vlastní varianta

^ Dítě

Příjmení	Doporučené	Vlastní
Holaubek	Holoubek (1)	Vlastní varianta

^ Otec

Jméno	Doporučené	Vlastní
Tobyass	Tobiáš (1)	Vlastní varianta

Příjmení	Doporučené	Vlastní
Hříbek	Žádné doporučené varianty	Vlastní varianta

Navrhnout varianty

Povolání	Doporučené	Vlastní
heytmán	Hejtman (1)	Vlastní varianta

Obec	Doporučené	Vlastní
Znoymenses	Znojmo (1)	Vlastní varianta

Obrázek 7. Ukázka normalizační části formuláře pro vytvoření/editaci záznamu ve webové aplikaci DEMoS. Formulář je rozdělen do sekcí. V prvním sloupci jsou vyplněny slova z původních polí formuláře. Druhý sloupec obsahuje doporučené normalizované varianty pro daná slova. Ve třetím sloupci je možné zadat vlastní normalizovanou variantu, pokud by žádná doporučená varianta uživateli nevyhovovala. Poslední sloupec obsahuje tlačítka pro explicitní navržení normalizovaných variant. Pro všechna slova, krom příjmení Hříbek, jsou dostupné normalizované varianty.

^ Porodní bába

Jméno	Doporučené	Vlastní
Terezija	✓ Terezie (2) Další Tereza (1) Nevybráno Nestandardizovat	Vlastní varianta

Obec	Doporučené	Vlastní
in Josefsthal		Vlastní varianta

Navrhnout varianty

Obrázek 8. Výběr doporučených normalizovaných variant z formulářového výběrového prvku. Pro dané jméno Terezija jsou na výběr dvě doporučené normalizované varianty. Slovo Terezija již bylo dvakrát normalizováno jako Terezie, proto varianta Terezie je jako první na výběr a je tedy i automaticky vybrána. Slovo Terezija se také nachází ve shluku slov, kde jedno z těchto slov již bylo normalizováno jako Tereza, proto v sekci Další je na výběr varianta Tereza. Poslední možností je nezvolit žádnou normalizovanou variantu.

Standardizace

Typ slov:	Jména	Slovo:	ate	Pohlaví:	Vše
Standardizováno:	Vše	Autor:	Jméno autora	Na stránce:	10
Filtrovat:	Filtrovat				

Slovo	Pohlaví	Standardizováno	Autor	Doporučené	Vlastní varianta	Potvrdit
Kateriny	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Kateriny	Ž	Catharina	7 počítač	Kateřina (1)		Potvrdit
Kateryna	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Katerzena	Ž	Catharina	7 počítač	Kateřina (4)	Katka	Potvrdit
Katerzina	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Katerziny	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Katerzyna	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Katerzyny	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit
Katerži	Ž	Catharina	7 počítač	Kateřina (4)		Potvrdit
Kateržina	Ž	Kateřina	7 DavidH	Kateřina (1)		Potvrdit

Obrázek 9. Stránka pro hromadnou normalizaci slov z matričních záznamů. V horní části se nachází filtrační formulář. Pod tímto fomulářem je tabulka pro normalizaci slov. Každý řádek tabulky odpovídá jednomu slovu k normalizaci. Ve sloupci Standardizováno jsou zobrazeny zvolené normalizované varianty pro jednotlivá slova (šedou barvou je zobrazeno číslo shluku tohoto slova). Ve sloupci autor je vypsán autor této normalizace. Dále je zde sloupec s doporučenými normalizovanými variantami a sloupec pro zápis vlastní varianty. Vlastní normalizovaná varianta má přednost před vybranou normalizovanou variantou ve sloupci doporučeno, na toto je uživatel upozorněn změnou barvy doporučených normalizovaných varianty při zadání vlastní varianty. Tlačítko potvrdit je možné stisknout pouze pokud je pro dané slovo vybrána normalizovaná varianta, která by vedla ke změně jeho normalizace.