

Digital humanities

Cvičení: Analýza textů, měření příbuznosti jazyků

Jindřich Marek

1. Analýza textů (frekvence, wordcloud)

- morfologická analýza, lemmatizace
- vizuální znázornění textových dat

Přehled

- toto cvičení se zaměřuje na analýzu textu a zpracování přirozeného jazyka (NLP) v českých textech
- naučíme se pracovat s morfologickou analýzou, lemmatizací, tvorbou wordclouds

Cíle cvičení

- po dokončení tohoto cvičení budete schopni:
 - pochopit základy NLP (zpracování přirozeného jazyka)
 - používat nástroj Voyant Tools pro vizualizaci textu
 - provádět lemmatizaci českých textů
 - vytvářet mračna slov (wordclouds)

Požadované nástroje

- webový prohlížeč pro Voyant Tools
- JupyterLab (jako v předchozích cvičeních)

Dataset

- „velké“ projevy českých prezidentů
 - vánoční poselství prezidenta Miloše Zemana 2022 (poslední)
 - novoroční projev prezidenta Petra Pavla 2024 (první)
- formát: plné texty v češtině

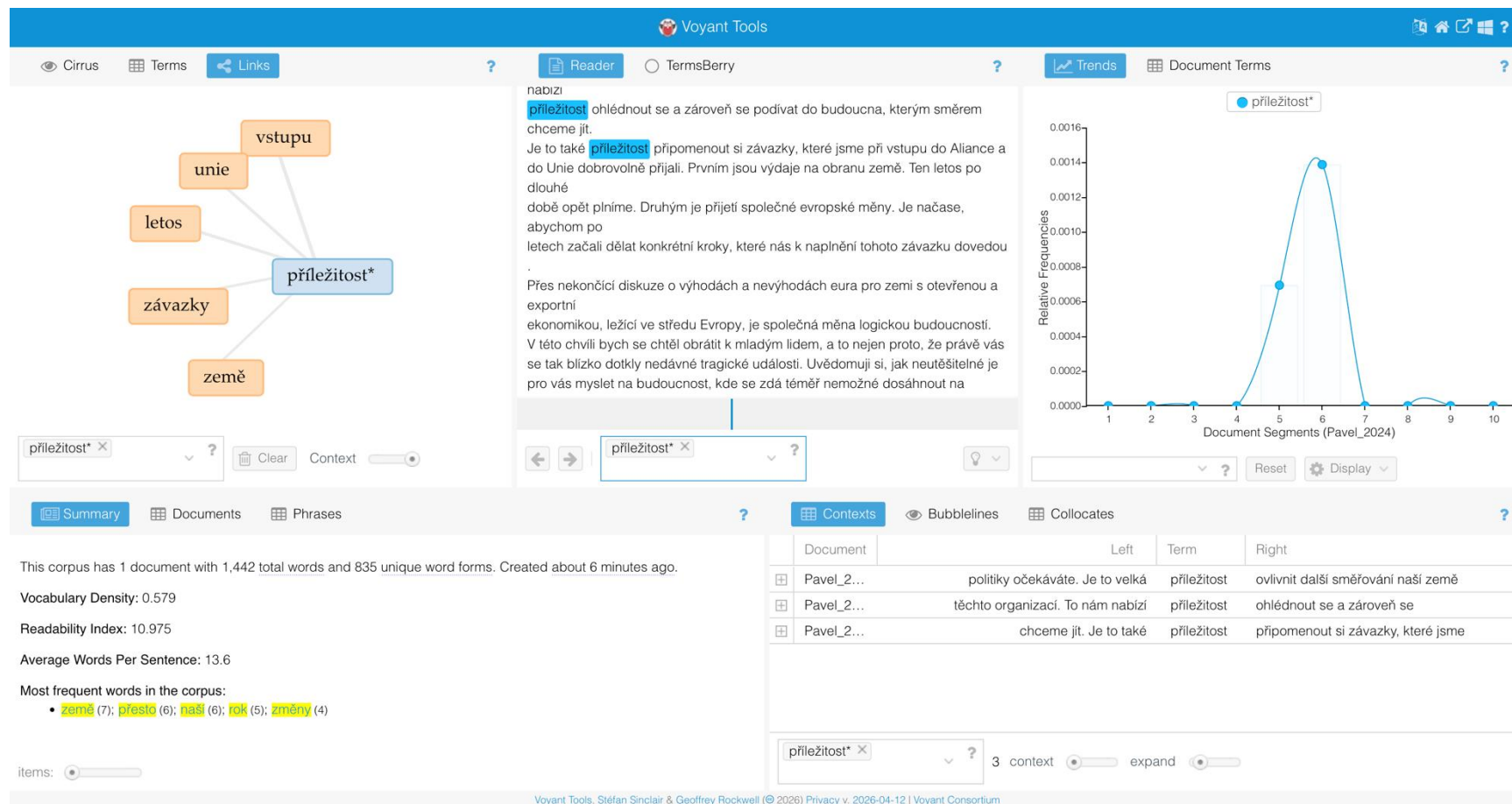
Postup

- porovnání slovní zásoby a témat

Rychlá vizualizace ve Voyant Tools

- <https://voyant-tools.org>
 - načtete soubor s projevem
 - nahrajte stop-slova (do uživatelem definovaného seznamu)
- prozkoumejte nástroje
 - Cirrus – slovní mrak
 - Terms – frekvenční seznam
 - Trends – rozložení slov v textu
 - Links – kolokace (která slova se vyskytují spolu)
- Která slova jsou nejfrekventovanější?
- Jak se liší oba projevy?
- Co Voyant Tools neukazuje?

Možný výsledek (Voyant Tools)



Notebook v prostředí Google Colab

The screenshot shows a Google Colab notebook interface. The top bar includes the notebook title, a star icon, and a share icon. Below this is a menu bar with options: Soubor, Upravit, Zobrazit, Vložit, Běh, Nástroje, and Nápoředa. A search bar labeled 'Přikazy' is on the left, and a toolbar with '+ Kód', '+ Text', and 'Spustit vše' is in the center. On the right, there are icons for chat, settings, a 'Sdílet' button, and a profile picture.

The left sidebar, titled 'Soubory', shows a file explorer with a folder named 'sample_data' and several files: 'Pavel_2024.png', 'Pavel_2024.txt', 'Zeman_2022.png', 'Zeman_2022.txt', and 'stopwords_cz.txt'. At the bottom of the sidebar, a 'Disk' section shows 'K dispozici: 86.44 GB'.

The main content area has a title 'Analýza dvou projevů prezidentů ČR'. Below the title, it says 'V tomto notebooku zpracujeme projevy dvou prezidentů:' followed by a bulleted list: 'Miloš Zeman (2022, poslední vánoční poselství)' and 'Petr Pavel (2024, první novoroční projev)'. Below the list, it says 'Na texty aplikujeme morfologickou analýzu pomocí knihovny Stanza, vyčistíme data a vytvoříme slovní mraky.'

The next section is titled 'Instalace knihoven'. It contains a paragraph explaining that the installation is problematic due to the large download of the torch library. It suggests running the notebook elsewhere, such as on a server (providing a URL: <https://hpc.troja.mff.cuni.cz:8000/>), or using a service like 'na e-infra' (noting that registration is required) or 'Google Colab' (noting that a Google account is needed). It also mentions that the notebook can be run locally.

The bottom section shows a code cell with the command `%pip install stanza wordcloud matplotlib pandas`. The output shows the installation progress for 'stanza', 'emoji', and 'udtools', including download sizes and requirement checks for various dependencies like numpy, platformdirs, protobuf, requests, networkx, torch, and tqdm.

The bottom status bar shows 'Proměnné', 'Terminál', a blue star icon, and '12:28 Python 3'.

Výsledek (Python, po lemmatizaci)

Miloš Zeman - vánoční poselství 2022



Petr Pavel – novoroční projev 2024



Limity frekvenční analýzy

- frekvence neznamená důležitost
 - častá slova nemusí být klíčová
- chybí kontext
 - nevíme, v jakém významu se slovo používá
- záleží na předzpracování
 - stop-slova, lemmatizace

Možná rozšíření

- kolokační analýza (jaká slova se vyskytují vedle sebe)
- analýza sentimentu (kladný vs. záporný tón)
- tematické modelování (LDA)
- porovnání v čase – jak se jazyk mění

Prompt pro jazykový model

Porovnej hlavní témata v plném textu následujících projevů.
Výstupem mají být stručné body na jeden snímek v prezentačním programu (levý a pravý sloupec).

Miloš Zeman (2022, poslední vánoční poselství)

[plný text]

Petr Pavel (2024, první novoroční projev)

[plný text]

Odpověď (Gemma 4 26B A4B)

MILOŠ ZEMAN (2022)

BILANCE A OBHAJOBA TRADICE

- **Aktuální krize:** Agrese Ruska (podpora Ukrajiny vs. ekonomické zájmy), energetická krize a Green Deal.
- **Ekonomická kritika:** Boj proti inflaci, kritika vysokých úrokových sazeb a deficitu státního rozpočtu.
- **Historický regres:** Přehled politické kariéry, obhajoba reform (samospráva, armáda, banky) a úspěchů sociální demokracie.
- **Národní zájem:** Obrana suverenity ČR, ekonomická diplomacie (export) a úzký vztah s Izraelem.
- **Osobní rozloučení:** Reflexe politického života, vyjádření vděčnosti a odchod do důchodu.

PETR PAVEL (2024)

VIZE A ODPOVĚDNOST ZA BUDOUCNOST

- **Bezpečnost a trauma:** Reakce na tragickou událost (střelba), překonání strachu a posílení komunitního bezpečí.
- **Politická kultura:** Kritika vulgarizace diskuse, apel na profesionalitu a vyhýbání se populistickým lákadlům.
- **Státní instituce:** Podpora koalice, řešení dluhu, důležitost opozice a aktivní role prezidenta v obraně práv.
- **Evropský směr:** Posílení role v EU a NATO, plnění závazků (obrana) a podpora přijetí společné měny (euro).
- **Generační odpovědnost:** Podpora mladých lidí, řešení jejich problémů (bydlení, duševní zdraví) a udržitelnost pro děti.

2. Měření příbuznosti jazyků

- lexikostatistika
 - metoda pro měření příbuznosti jazyků na základě slovní zásoby.
- princip
 - vybereme základní slovní zásobu
 - Swadeshův seznam – slova jako já, ty, hlava, voda, oheň
 - porovnáme podobnost slov mezi jazyky
 - vytvoříme vizualizaci
 - např. dendrogram (strom podobnosti)

Přehled

- toto cvičení se zaměřuje na analýzu podobností mezi jazyky

Cíle cvičení

- po dokončení tohoto cvičení budete schopni:
 - měřit podobnost mezi jazyky pomocí lexikálních dat

Požadované nástroje

- JupyterLab (jako v předchozích cvičeních)

Dataset

- přehled domorodých jazyků jižní Ameriky
 - formát: tabulka se slovy specifickými pro jednotlivé jazyky

Postup

- seznam vybraných slov
- očištění slov na základní tvary
- analýza podobnosti slov
 - hierarchické shlukování
- vizualizace
 - dendrogram
 - matice podobnosti
 - zobrazení ve dvourozměrném prostoru

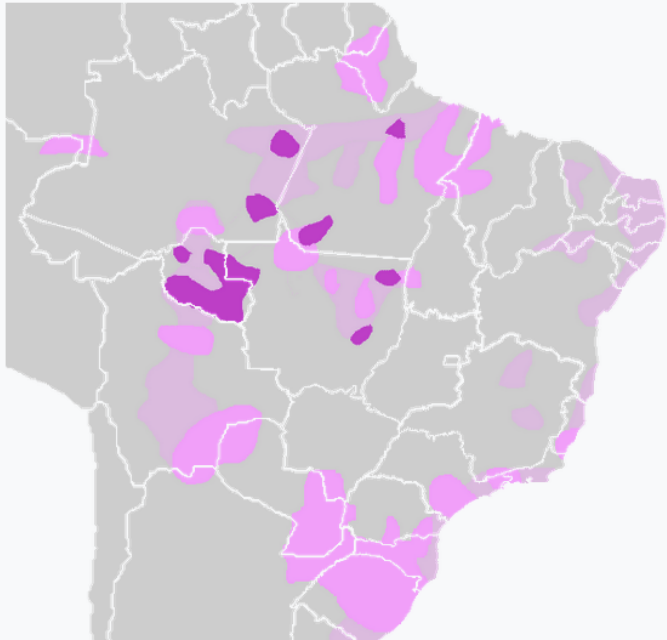
Tupíjské jazyky neboli *jazyková rodina tupí* jsou **rodinou** jazyků, kterými mluví domorodí obyvatelé **Jižní Ameriky**, konkrétně v **Amazonii**, na východním pobřeží **Brazílie**, v **Peru**, **Bolívii**, **Paraguayi** a **Francouzské Guyaně**.

Do rodiny se řadí okolo 70 jazyků. Nejvýznamnější skupinou tupíjské rodiny je skupina **tupí-guaraní**, kam patří **guaranština**, jíž mluví okolo 7 miliónů lidí a v Paraguayi a Bolívii požívá statusu **úředního jazyka**.^[1] Dalšími tupíjskými jazyky mluví dnes daleko méně lidí: jazyk *cocama* v Peru má 19 000 mluvčích, jazykem *chiriguano* v Bolívii hovoří přibližně 20 000 lidí; ostatní jazyky mají tisíce či stovky mluvčích.^[2]

V minulosti se některé tupíjské jazyky (respektive jejich zjednodušené formy) používaly v Brazílii jako **lingua franca** (všeobecný dorozumívací prostředek) mezi misionáři a domorodým obyvatelstvem.^[1]

Jeden z vymřelých jazyků, *tupinambština*, je zdrojem mnoha slov z oblasti pojmenování jihoamerické fauny a flóry v portugalštině, odkud se tato slova šířila i do dalších evropských jazyků.^[3] Z těchto slov v češtině používáme třeba slova **jaguár** či **tamarín**.^[4]

Tupíjské jazyky



Území, kde se mluví či mluvilo tupíjskými jazyky

Rozšíření	Jižní Amerika
Počet mluvčích	není znám
Počet jazyků	70
ISO 639-2	tup

Život [[editovat](#) | [editovat zdroj](#)]

Jeho největší přínos spočívá v dodnes užívané klasifikaci indiánských **jazyků Jižní Ameriky** na lexikálním základě. Metody klasifikace jazyků uplatnil také pro papuánské jazyky. Vydal knihu o vývoji písma.^[3]

Dílo [[editovat](#) | [editovat zdroj](#)]

- *Roztřídění jihoamerických jazyků*, Praha, 1935, též španělsky jako *Clasificación de las lenguas sudamericanas*, Praha, 1935 a posmrtně anglicky jako *Classification of South American Indian Languages*, Berkeley, 1968.
- *Vývoj písma*, Praha, 1946.
- *Classification des langues papoues*, *Lingua Posnaniensis* 6 (Poznań), s. 19–83.

A black and white portrait of a middle-aged man with thinning hair, wearing round-rimmed glasses and a dark suit jacket over a light-colored shirt. He is looking directly at the camera with a neutral expression. The background is out of focus, showing some foliage.

Čestmír Loukotka (Pražský ilustrovaný
zpravodaj, 1941)

Narození 12. listopadu 1895
Chrástany

 Rakousko-Uhersko

Úmrtí 13. dubna 1966 (ve věku 70 let)
Praha
 Československo

Povolání jazykovědec, antropolog a etnograf

Děti

- Jarmila Loukotková dcera

 **multimediální obsah** na Commons

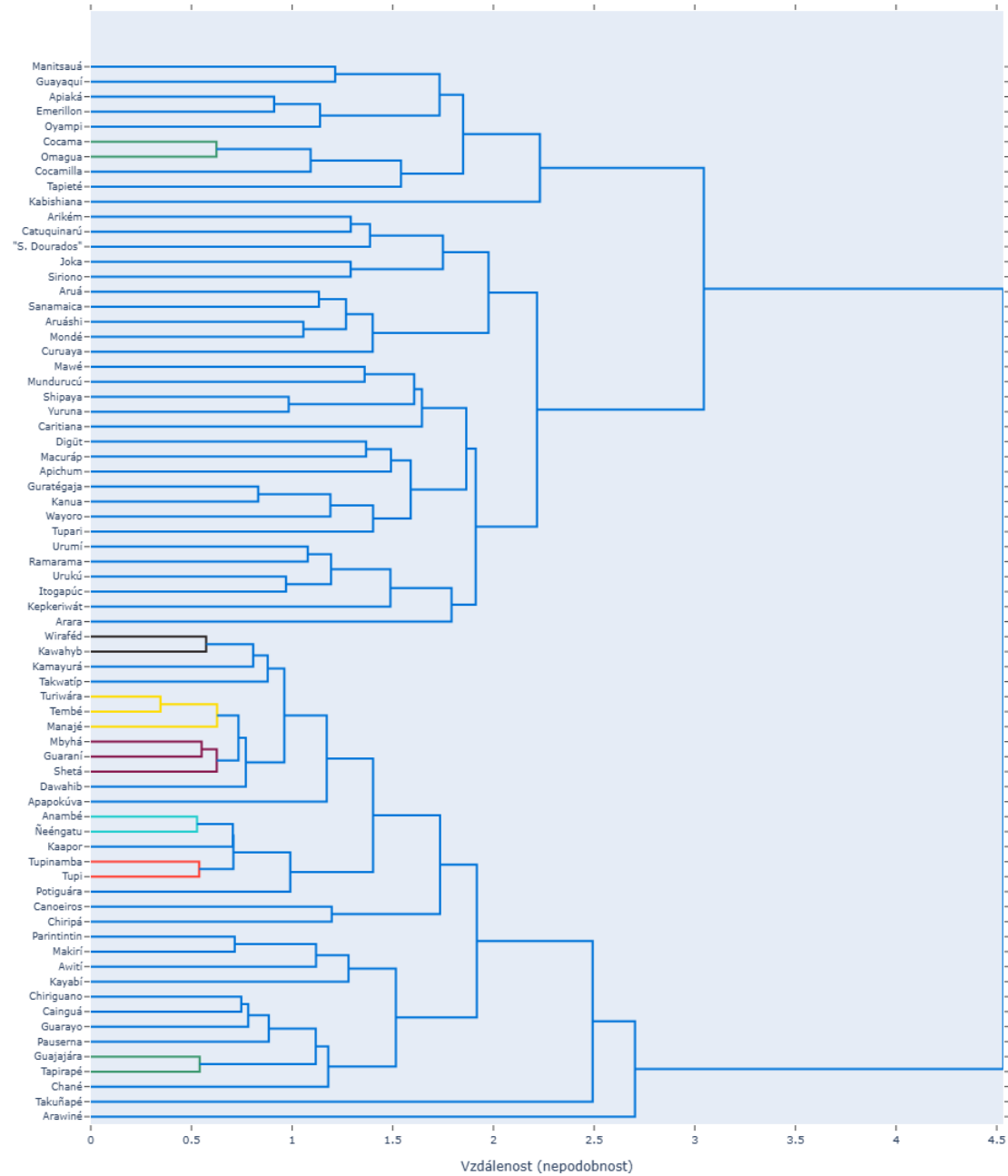
Některá data mohou pocházet z datové položky.

Vocabulary [\[edit \]](#)

Loukotka (1968) lists the following basic vocabulary items.^[6]

Language ♠	Branch ♠	head ♠	ear ♠	tooth ♠	hand ♠	one ♠	two ♠	three ♠	woman ♠	water ♠	fire ♠	stone ♠	maize ♠	tapir ♠
Tupi	Tupi	a-kang	nambi	táña	pó	peteĩ	mokoĩ	mbohapüi	kuñá	ü	tatá	itá	abai	tapũíra
Tupinamba	Tupi	a-kán	nambü	ráña	pó	angepé	mokoin	musaput	kuñá	ü	tatá	itá	auvati	tapirusu
Potiguára	Tupi	a-kanga	nambi	tanha	in-bó	oyepe	mokoy	mosapür	kuña	üü	tata:	ita:		
Ñéengatu	Tupi	a-kanga	namü	taña	pó	yepé	mokoin	musapeire	kuñan	üüg	tatá	itá	auati	tapira
Guaraní	Guaraní	ãkan	nambi	apen-kun	pó	peteĩ	mokõi	mbohapy	kuñá	ü	tatá	itá	avatí	tapií
Apapokúva	Guaraní				pó	aépi	mokõi	moapi	kuña	ü	tatá			
Chiripá	Guaraní	rakã	nambi			aépi				ü	tata		avati	mborevi
Caingúa	Guaraní	aká	nambi	kú	pó	petein	mókoin	mbohapi	koñá	ü	tatá	itá	avachi	mborevi
Mbyá	Guaraní	che-ahká	chen-nambüh	che-rain	cheh-pó	peteí	mokoi	mboapü	kuña	ü	tatá	itá	avachi	tapií
Canoeiros	Guaraní	eaushmä			de-pó				uainvi	üg		itá	avashi	
Shetá	Guaranized	sh-aka	che-nambi	tienai	che-pó	matinkam	mokoi	ñiiru	kuñá	ü	tatã	itá	avachi	tapi
"S. Dourados"	Guaranized	ñ-ãka	ela:me	nénai	e:-po	uazi	mo:gai	mágatei	ko:ña	ho:ñe	agel'á	i:tá	nutya	tela:goi
Guayaquí	Guaranized	ni-aka	nambi	ã	i-pá	eteyã	meno	tanã	kuña	ü	dadá	itá	waté	mberevi
Tapirapé	Tapirapé	dzyane-akánga	dzyane-inamí	dzyane-roi	dzyane-pó	anchepe	mukúi	mãpít	kudzá	ü	tatá	itá	awachí	tapiíra
Kamayurá	Kamayurá	ye-akang	ye-nami	ye-nai	ye-po	yepete	mokoi	moapit	kuña	ü	tata	ita	avatsi	tapiít
Awití	Kamayurá	apot	inte-yambe	inte-ngu	i-po	mayepete	monkói	munitaruka	kuñá	ü	tara	ita	avachi	tapií
Arawiné	Kamayurá		ne-nami		ye-po									
Anambé	Pará	a-kánga	hã-nambi	se-raña	pó	yanäpo	mukuẽ	muhapi	kuña	ü	tata	ita	awat	tapiri

Hierarchické shlukování tupijských jazyků na základě lexikálních podobností (Plotly)



Zobrazit popisky Skrýt popisky



Limity měření podobnosti

- měříme podobnost znaků – co když se slova liší jen v jednom znaku, ale mají stejný původ?
- ignorujeme sémantiku – co když se význam slova v čase změnil?
- malý vzorek slov – stačí 13 slov pro spolehlivou klasifikaci?

Klíčové poznatky a otázky k diskusi

- klíčové poznatky
 - tupijské jazyky tvoří jasně oddělitelné větve
 - podobnost slovní zásoby koreluje s klasifikací lingvistů, ale metoda má své limity
- co ovlivňuje podobnost slovní zásoby kromě příbuznosti?
 - geografická blízkost
 - jazykový kontakt
 - doba oddělení
- možná rozšíření
 - použít větší Swadeshův seznam (100 nebo 200 slov)
 - zahrnout fonologickou podobnost
 - porovnat s výsledky komparativní metody