

Korpusová lingvistika

- Studie jazyka:
- struktury
 - užití

Co to je korpus? Sbírkou textů, v ideálním případě kompletní
- je anotovaný nějak
- Kladný korpus kade přirozeně získaný.
Něco se neobjeví, kvůli nějakým problémům.

Vlastnosti:

- reprezentativnost
- končinná velikost
- strojově čitelná forma
- standardní reference

→ Je důležité v jistý moment zastavit. Protože
je potřeba porovnat výskyt v jiných časích.

Wáňkové korpusy

Brown korpus

- první velký korpus
- snažil se nejlépe reprezentovat americkou angličtinu v roce 1961.

Penn Treebank

- první syntakticky anotovaný korpus
- články z Wall Street Journal → Je to ale burzovní slang, není to spisovná angličtina.

Český národní korpus

- vzniklo v 90. letech : UK, MV, AU
 - anotáce na morfologické úrovni
 - mix novin, knih a technické literatury
- 60% → 25% → 15%

Pražský závislostní korpus (PDT) (100 000 vět, 1 200 000 slov)

- syntakticky anotovaný korpus → celý anotovaný ručně (na rozdíl od ostatních ve světě)
- založeno na teorii funkčního generativního popisu
- úrovně anotace:
 - morfologie
 - analytická rovina (pouze syntaktická)
 - funkční gramatická rovina: závislostní, jicho, koreference, vše ostatní

Prague Ambie Dependency Treebank

- první velký závislostní korpus použitý na typově úplně odlišném jazyku.

Pravděpodobnostní a statistické metody

- pravděpodobnost je jen jen aproximační trénovací dat
- $\hookrightarrow$ tam může být velmi odlišná
- například trénujeme překlad na paralelním slovníku / korpusu.

Bayesův vzorec:
$$P(A|B) = \frac{P(B|A) \cdot P(A)}{P(B)}$$

- Cíl: chceme předpovědět $p(w|h)$, kde w je hledané slovo a h je historie textu.
- tyto modely by byly ale obří a většinou přehradí.

Používám se trigramový model:

jako „okénko“ historie bude pracovat kontext délky 3.

- spousta konstantní bude mít 0 pravděpodobnost

Metoda vyhlazování

- nulové pravděpodobnosti se nahradí na velmi malý ϵ .

- už ale nerozumně neexistující kombinace

- Příklad boletné není dobrý trénovací dataset, jelikož se preferuje zachování významu před formou 1:1.

- Do 2015 se předkládalo statisticky, nyní funguje neuronový předklad.

Metoda záměrného klanění: chceme $f: F \rightarrow A$, hledáme $P(A|F)$

Bayes mi říká:
$$P(A|F) = \frac{P(F|A) \cdot P(A)}{P(F)}$$

$\rightarrow$ tady můžeme chápat jako $= 1$
když jsme to dostali na vstup

problém ztráty
negací

postřední sémantické
pochopení

modely: předkladový $P(f|a)$

cílového jazyka $P(a)$

- těch hypotéz musíme být ještě mnoho?

$\hookrightarrow$ obrátíme směr předkladu a hledáme větu, ze které to vzniklo.

- neexistuje vztah k originálu, kde se vyberecí cílová věta

- předklad problém obecně

BLEU:

$$\text{BLEU: } BP * (p_1 p_2 p_3 p_n)^{1/n} \in (0, 1)$$

→ lze reprezentovat jako %, ale není to určeno pro porovnávání odlišných systémů

↑
penalizace za stručnost

↓
geometrický průměr n-gramové přesnosti pro $n=1 \dots n$

⇓
což je porovnávací délka,
aby zvládla kompletní větu

$$p_i := \# \text{ výskytů } i\text{-gramů} / i\text{-gramů celkem}$$

- je to velmi rychlé
- velmi přesné, pokud existuje více ref. překladů
 - použití synonyma by muselo škodit, kdybych měl jen jeden ref. překlad
- adjoinování trénovacího datasetu zvedne BLEU pouze o 0,5%
 - nemůžeme jen tak stáhnout celý internet, jelikož spousta stránek je vytvořeno již výplňovým strojovým překladem, takže se to nedá použít.