

[Inspiratron.org](https://inspiratron.org) - Natural language processing, machine learning and cybersecurity

Klasifikator sentimenta za srpski jezik

by Nikola Milošević - Saturday, October 20, 2012

<https://inspiratron.org/blog/2012/10/20/klasifikator-sentimenta-za-srpski-jezik/>

S obzirom da ulazim u poslednju fazu mog master rada, koji se bavi mašinskom sentiment analizom srpskog jezika, mislim da bi bilo dobro da napišem par reči o tome ovde. Sentiment analizator je zapravo i razlog zakupa ovog domena. Evo o čemu se radi:

Analiza sentimenta je oblast mašinskog procesuiranja prirodnog čovekovog govora čiji je cilj otkriti sentiment određenog teksta ili rečenice, odnosno subjektivan osećaj tog dela teksta. Najprostije rečeno da li rečenica ima pozitivan ili negativan kontekst. Radi se o klasifikaciji teksta u dve klase u ovom slučaju - pozitivnu i negativnu. Klasifikacija može biti i u više klasa (mnogo pozitivno, pozitivno, neutralo, negativno, mnogo negativno i sl.), međutim ovde ćemo se baviti samo binarnom klasifikacijom. Prilikom ove klasifikacije postoje određeni izazovi, koji je razlikuju od obične klasifikacije teksta po temama. Recimo, negacija utiče mnogo više na sentiment rečenice nego na njegovu temu. Pa tako rečenica "Ovo je dobar čovek" i "Ovo nije dobar čovek" po klasifikaciji teme su jako slični, ali mnogo različiti kada se radi o sentimentu. Naravno, postoje određeni problemi, koji još u naučnim krugovima su predmet istraživanja i koji nisu adekvatno rešeni, poput ironije.

Klasifikacija teksta i analiza sentimenta su naročito bitni u mnogim sociološkim i poslovnim procesima. Tema je postala naročito interesantna nakon 2001. godine, kada je ekspanzija interneta, blogova, foruma i socijalnih mreža donela velike mogućnosti za eksploataciju analize sentimenta i klasifikaciju uopšte. Danas, analiza sentimenta je bitna kako u poslovnoj sferi, gde se želi saznati utisak o određenom proizvodu na osnovu podataka prikupljenih sa interneta, preko politike gde se želi saznati kakav je utisak javnosti o određenom kandidatu, do socijologije koja može imati veliki benefit od razvoja ove oblasti. Takođe klasifikacija se koristi u različitim servisima za online oglašavanje, često i zajedno sa sentiment analizom, jer ukoliko je tekst o određenom proizvodu napisan u negativnom kontekstu ne bi bilo dobro pored teksta postaviti reklamu sa baš tim proizvodom. Naravno navedeni načini korišćenja su mnogo brži i jeftiniji od klasičnog istraživanja javnosti ili pak donose dodatnu vrednost u oglašavanju.

Za određeni klasifikator je iskorišćen Naive Bayes algoritam i posebno razvijen stemer za srpski jezik. Prema mojim informacijama radi se o jednom od prvih sentiment analizatora za srpski jezik, kao i jedan od prvih uređenih stemera za srpski jezik sa malim brojem pravila (oko

300, raniji stemer za koji znam je stemer Danka Šipke i Vlade Kešnja sa preko 1000 pravila i manjom tačnošću, bar za task analize sentimenta).

O stemeru ću pisati posebno, tako da ću sad objasniti kako radi Naive Bayes klasifikator.

Naive Bayes je algoritam supervizovanog mašinskog učenja, što znači da postoji trening set podataka koji su labelisani i gde su određene klase dokumenata. Na osnovu ovih dokumenata i reči u njima se radi dalja statistička analiza i određuje klasa novih, nepoznatih dokumenata.

Naive Bayes klasifikator je zasnovan na Bayesovom pravilu koje glasi:

$$p(C|F_1, \dots, F_n) = \frac{p(C) p(F_1, \dots, F_n|C)}{p(F_1, \dots, F_n)}.$$

Gde je C klasa, a F feature, odnosno reč. Interpretacija glasi da je verovatnoća da reči F1 do Fn pripadaju klasi C, jednaka verovatnoći klase C pomnožene sa verovatnoćom da se u klasi C nalaze feature-i F1 do Fn, podeljenom sa verovatnoćom feature-a. U praksi interesuje nas samo gornji deo, jer donji ne zavisi od C. Pa se tako jednačina može napisati kao:

$$\begin{aligned} p(C, F_1, \dots, F_n) & \\ & \propto p(C) p(F_1, \dots, F_n|C) \\ & \propto p(C) p(F_1|C) p(F_2, \dots, F_n|C, F_1) \\ & \propto p(C) p(F_1|C) p(F_2|C, F_1) p(F_3, \dots, F_n|C, F_1, F_2) \\ & \propto p(C) p(F_1|C) p(F_2|C, F_1) p(F_3|C, F_1, F_2) p(F_4, \dots, F_n|C, F_1, F_2, F_3) \\ & \propto p(C) p(F_1|C) p(F_2|C, F_1) p(F_3|C, F_1, F_2) \dots p(F_n|C, F_1, F_2, F_3, \dots, F_{n-1}). \end{aligned}$$

Ovde su uvodi naivni pretpostavka da feature-i, odnosno reči F1...Fn ne zavise jedni od drugih, odnosno da su potpuno nezavisni. Tako jednačina dobija oblik:

$$\begin{aligned} p(C, F_1, \dots, F_n) & \propto p(C) p(F_1|C) p(F_2|C) p(F_3|C) \dots \\ & \propto p(C) \prod_{i=1}^n p(F_i|C). \end{aligned}$$

Odnosno:

$$p(C|F_1, \dots, F_n) = \frac{1}{Z} p(C) \prod_{i=1}^n p(F_i|C)$$

Odnosno najverovatnija klasa je ona čija je verovatnoća najveća. Računa se nad istim rečima i za pozitivnu i za negativnu, uporede, i klasa je ona čija je verovatnoća veća.

Verovatnoća klase se računa kao zbir svih dokumenata određene klase, podeljena sa zbirom svih dokumenata. Verovatnoća reči u klasi se računa kao zbir pojavljivanja određene reči u toj klasi podeljen sa zbirom pojavljivanja svih reči u trening setu.

Problem u ovakvom izražavanju postoji sa rečima koje ne postoje u trening setu. Njihova verovatnoća bi u ovom slučaju bila 0, što bi poremetilo ceo proces određivanja klase. Zato se uvodi poravnavanje koje predpostavlja da su i ti dokumenti javljeni jednom, kao i da svi ostali dokumenti su se javili za jedan više put nego što zaista jesu. Ovakvo poravnavanje se zove Laplasovo poravnavanje (Laplace smoothing). Na ovaj način je moguće dobiti kvalitetne klase i za nove podatke koji se nisu javili u trening setu.

Pre ubacivanja teksta u Naive Bayes klasifikator, bilo da se radi o učenju ili da se radi o određivanju sentimenta novih rečenica potrebno je uraditi nekoliko obrada. Prvo je potrebno uraditi stemming, odnosno uklanjanje sufiksa, kako bi reči različitih fleksija imali isti oblik. Potrebno je takođe dovesti sva slova na istu veličinu, odnosno smanjiti velika slova da budu mala. Takođe urađena je i obrada negacija. Nakon pojavljivanja reči ne ili glagolske negacije glagola jesam (nije) ili hteti (neću), dodat je prefix NE_ svim rečima do znaka interpunkcije. Na taj način je obezbeđeno da te reči iz rečenice negativnog konteksta ne budu zajedno sabirani sa rečenicama pozitivnog konteksta u bazi.

Detaljnije o svemu ovome, strukturi baze i samom algoritmu, kada odbranim rad.

Nikola Milošević

All rights reserved and copyrighted by inspiratron.org and Nikola Milosevic