

Abstract

Information retrieval systems are very important as they permit us to easily access huge amounts of data. In the current context where we are having a lot of information, the precision of an information retrieval system is crucial. Besides retrieving content, having mechanisms that can compare different entities of text and group them by similarity is extremely important. A model based on time series can be used as a viable alternative to the well-known vector space model. The time series model takes into account the order in which the words appear and moreover it can be easily extended, in order to add semantics by defining a semantic function on individual words. The generic model that we are proposing can be easily integrated with different criteria in order to compute distances between text corpuses and thus to find text entities that are similar to each other. Given the unstructured nature of the natural language, this problem raises different challenges in terms of both precision (it is difficult to determine the similarity of two or more entities in the absence of structure) and performance (analyzing unstructured data requires complex algorithms that usually have a high time complexity). Therefore, the approach proposed in this work provides: a generic framework that makes use of time series, which can be used in information retrieval and natural language processing, improves the similar text search results quality and identifies texts that belongs to the same source.

Sumar

Sistemele de regăsire a informației sunt foarte importante, deoarece cu ajutorul lor putem accesa cantități mari de date. În momentul de față, avem acces la foarte multe date, motiv pentru care precizia unui astfel de sistem este esențială. De asemenea, este important ca pe lângă găsirea unor informații să dispunem de un sistem capabil să compare două entități de text și să le grupeze pe baza similitudinii. Astfel, un model bazat pe serii de timp poate reprezenta o alternativă la bine cunoscutul model bazat pe spații vectoriale. După cum vom vedea, modelul bazat pe serii de timp ține cont de ordinea de apariție a cuvintelor în propoziție și poate fi extins ușor, pentru a ține cont și de semantică – acest lucru se realizează ușor prin integrarea unei funcții semantice de calcul a distanței dintre cuvinte. Modelul generic pe care îl propunem poate fi ușor extins cu alte criterii care pot fi folosite pentru a calcula distanța și, respectiv, similitudinea dintre texte. Natura nestructurată a limbajului natural duce la apariția diferitelor provocări, datorită cerințelor de precizie (este dificil de determinat similitudinea dintre două sau mai multe documente în absența structurii) și performanța (analiza datelor nestructurate necesită algoritmi care, în general, au complexitate ridicată). Astfel, această lucrare propune: o arhitectură bazată pe serii de timp, care poate fi folosită atât pentru regăsirea informațiilor, cât și pentru prelucrarea limbajului natural, o metoda care îmbunătățește rezultatele obținute la calcularea textelor similare și o metodă care grupează textele pe baza autorului.