



Міжнародна науково-технічна конференція

**ІНТЕЛЕКТУАЛЬНІ ТЕХНОЛОГІЇ ЛІНГВІСТИЧНОГО
АНАЛІЗУ**

23-24 жовтня 2013 року

ПІДХОДИ ДО АВТОМАТИЧНОГО ВИЗНАЧЕННЯ ТЕРМІНОЛОГІЧНИХ ОСНОВ ОНТОЛОГІЙ

Д.В. ЛАНДЕ, д.т.н.,

завідувач відділом

**Інституту проблем реєстрації інформації НАН
України**

Київ, 23 жовтня 2013 р.



Анотація

Доповідь присвячено актуальному на цей час питанню визначення важливих структурних елементів тексту, що виявляються інформаційно-значущими, визначають інформаційну структуру. Використання таких елементів дозволяє формувати тезауруси, пошукові образи документів, онтології.



Існуючі підходи

Ключові слова для пошуку в тексті, опорні слова для автоматичного екстрагування значущих фрагментів текстів або формування автоматичних рефератів, вибираються з урахуванням такої властивості слів, як «дискримінантна сила», для визначення яких існує три класичних методи – TFIDF, дисперсійний, ентропійний. Також застосовуються й новітні мережеві методи.



Показник TFIDF

Для кожного терміну з текстового корпусу у канонічному виді TFIDF дорівнює добутку частоти слова у фрагменті тексту (Term Frequency) на логарифм від величини, зворотної кількості окремих фрагментів тексту, в яких це слово зустрілось (Inverse Document Frequency). **СЕРЕДНЄ ЗА МАСИВОМ.** Для кожного слова i , що входить до текстового корпусу, який складається з N документів, підраховується кількість документів $df(i)$, в яких міститься це слово, а також загальна частота входження даного слова у текстовий корпус - $n(i)$. Після цього розраховується середнє значення TFIDF для кожного слова:

$$tfidf(i) = \frac{n(i)}{N} \log \left(\frac{N}{df(i)} \right)$$



TFIDF Окарі BM25

TFIDF дає не зовсім коректні результати на текстових масивах із документів з різною за довжиною. Тому запропоновано його модифікацію – Окарі BM25:

$$TF \times IDF(w_i) = IDF(w_i) \frac{f(w_i, D) \times (k_1 + 1)}{f(w_i, D) + k_1 \times (1 - b + b \times \frac{|D|}{avgdl})}$$

де $f(w_i, D)$ – частота слова у документі D ; $|D|$ – довжина документа D ; $avgdl$ – середня довжина документа в колекції; k_1 і b – параметри, що вибираються експертно. IDF обчислюється за формулою:

$$IDF(w_i) = \log \frac{N - n(w_i) + 0.5}{n(w_i) + 0.5},$$

де N – загальна кількість документів у масиві, $n(w_i)$ – кількість документів, що містять термін w_i .



Дисперсійна оцінка

Дисперсійна оцінка важливості термінів обчислюється для окремих термінів наступним чином. Позначимо середній інтервал між появами слова A у тексті через $\langle \Delta A \rangle$, а середній квадрат значень цих інтервалів через $\langle \Delta A^2 \rangle$. Дисперсійна оцінка слова σ_A розраховується як

$$\sigma_A = \frac{\sqrt{\langle \Delta A^2 \rangle - \langle \Delta A \rangle^2}}{\langle \Delta A \rangle}.$$

Ідея дисперсійної оцінки близька до TFIDF, однак більш коректно застосовується до цілих текстів, а не до масивів з великої кількості документів, як TFIDF.



Оцінка PMI

Якщо розглядати терміни як події, можемо знайти так звану ентропійну оцінку ваги терміну a (Pointwise Mutual Information, PMI):

$$m(w) = -\log_2 \frac{p(w | d)}{p(w)} = -\log_2 \frac{n_{w,d}}{n_w},$$

де $n_{w,d}$ – кількість входжень терміну до тексту документа, n_w – загальна кількість входжень.



Мережі мови

Першим кроком при застосуванні теорії складних мереж до аналізу тексту є представлення цього тексту у вигляді сукупності вузлів і зв'язків, побудова мережі мови (**Language Network**).

Існують різні інтерпретації значень вузлів і зв'язків, що веде, відповідно до різних представлень мережі мови.

Вузли можуть бути поєднані між собою, якщо відповідні їм слова знаходяться поруч у тексті, належать одному реченню, поєднані синтаксично або семантично.



Простіші мережі мови

L-простір. Зв'язуються сусідні слова, що належать одному реченню. Кількість сусідів для кожного слова (вікно слова) визначається радіусом взаємодії R , найчастіше розглядається випадок $R = 1$.

V-простір. Розглядаються вузли двох типів, що відповідають реченням і словам, які входять до них.

P-простір. Всі слова, що належать одному реченню, зв'язуються між собою.

S-простір. Речення зв'язуються між собою, якщо у них застосовуються однакові слова.

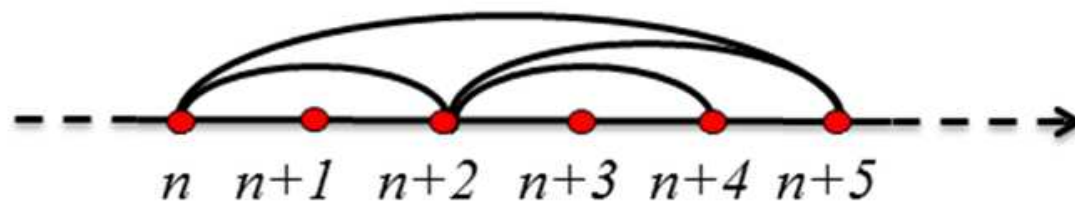
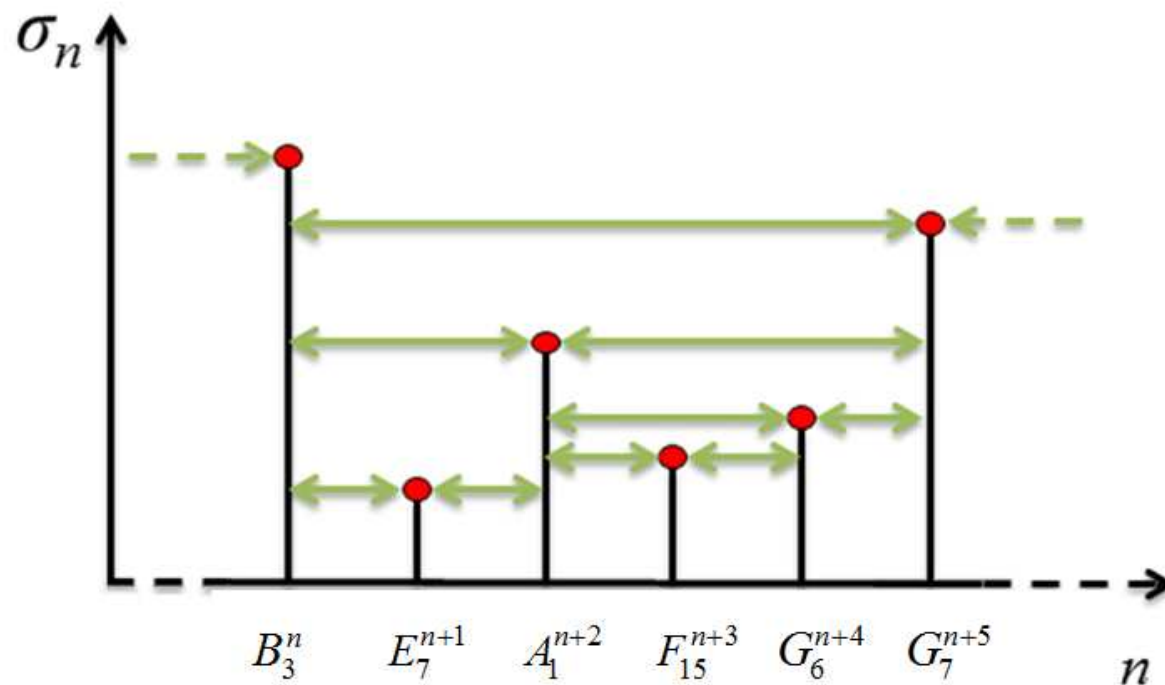


Побудова графів на основі рядів значень

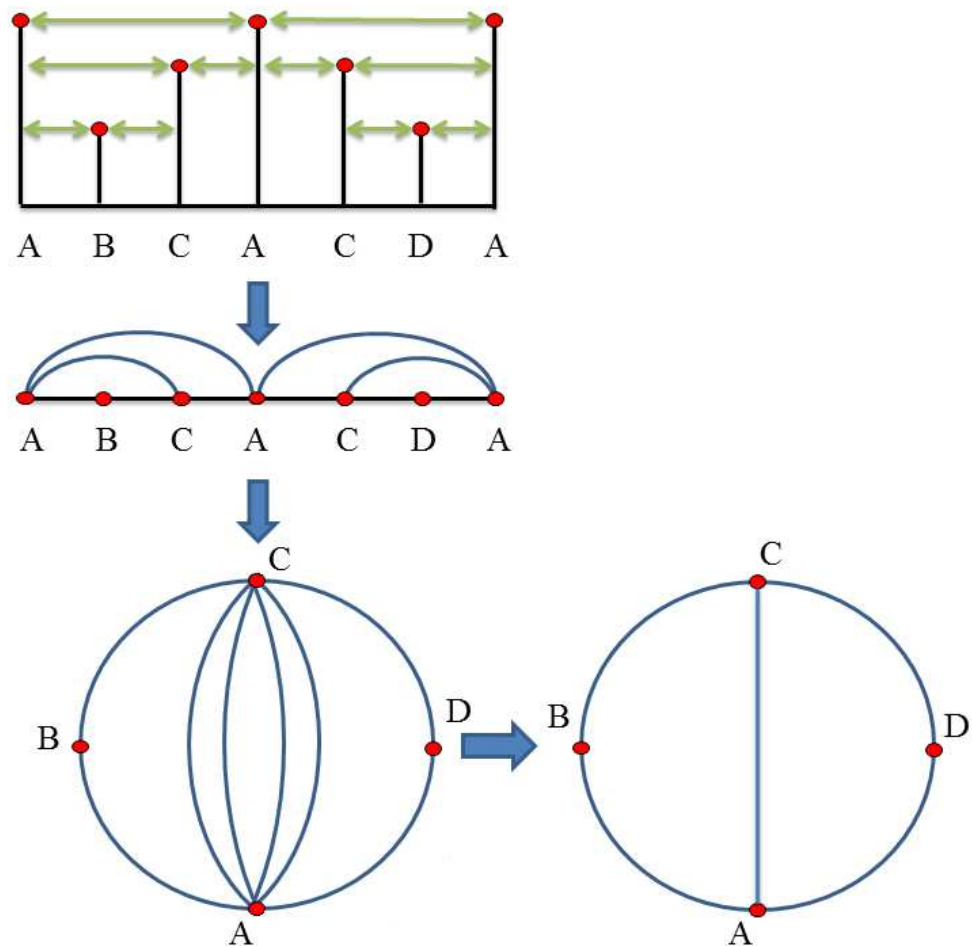
На цей час запропоновано декілька методів побудови мереж на основі часових рядів, серед яких можна назвати методи побудови графів видимості, зокрема так званий граф горизонтальної видимості (Horizontal Visibility Graph – HVG). Ці підходи дозволяють будувати мережеві структури також і на основі текстів, у яких окремим словам або словосполученням деяким чином поставлені у відповідність деякі вагові значення.



Формування графа горизонтальної видимості

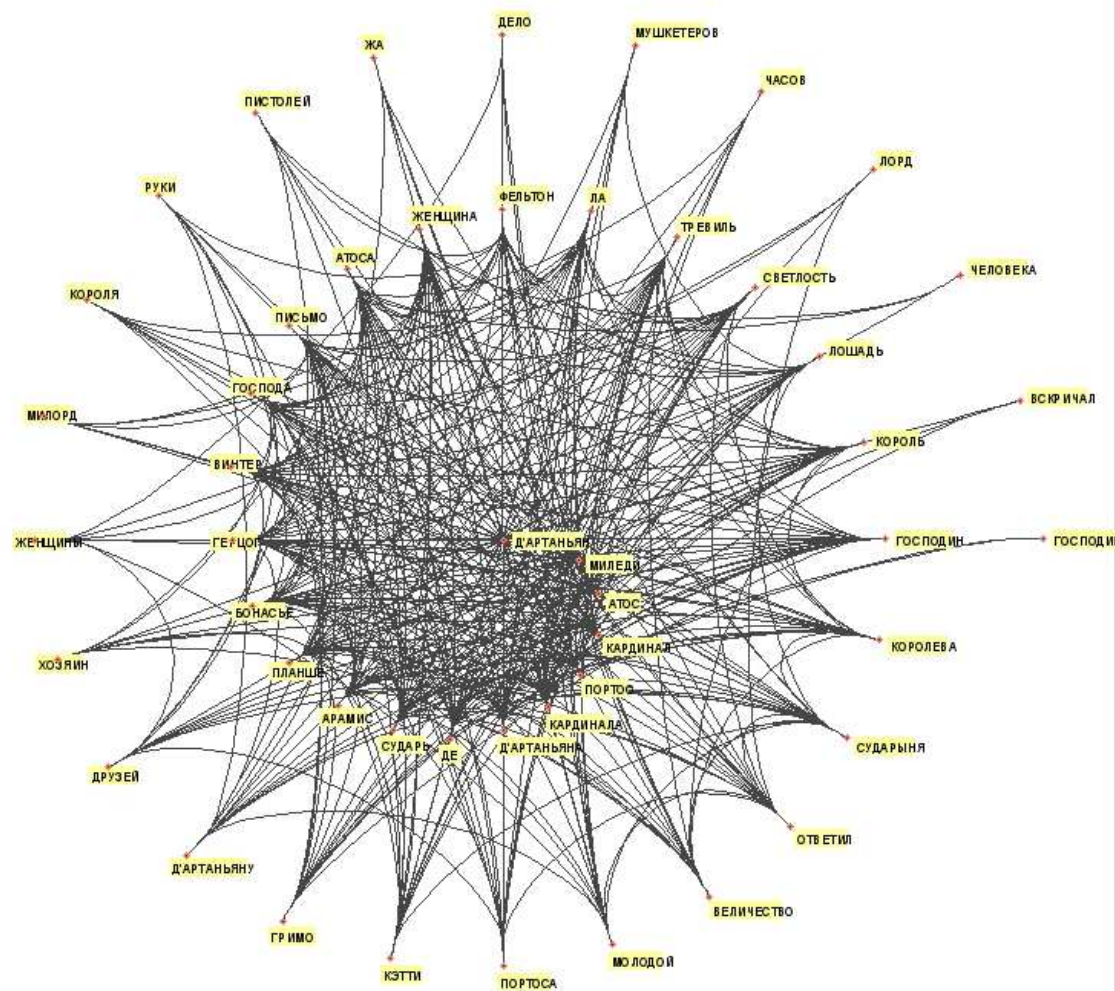


Етапи побудови КГГВ

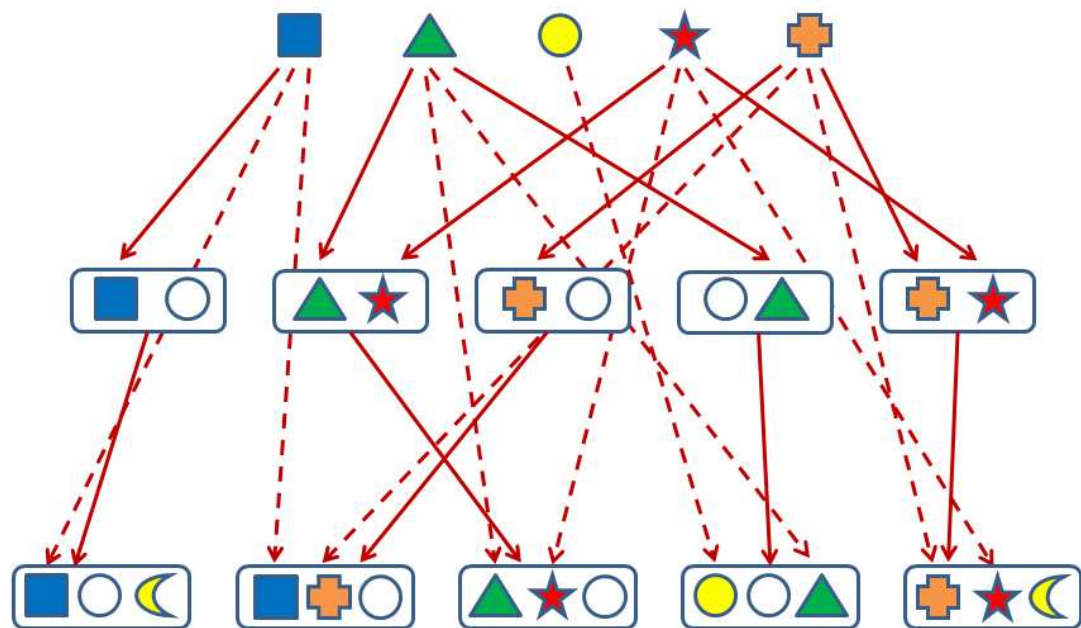


Компактифікований граф
горизонтальної видимості

О. Дюма, «Три мушкетера»



Мережі природних ієрархій термінів



З відібраних термінів (уніграм, біграм, триграм) формуються мережі природних ієрархій термінів, у яких як вузли розглядаються самі терміни, а зв'язкам відповідають входження одних термінів у інші.



Вихідний текстовий корпус

Як приклад вихідного текстового корпусу розглядаються анотації електронних препринтів Arxiv (<http://arxiv.org>) за тематикою інформаційного пошуку (рубрика cs.IR, понад 500 документів за 5 років).

The screenshot shows the Cornell University Library's Arxiv.org interface. At the top, the Cornell University Library logo is on the left, and a note about support from the Simons Foundation is on the right. Below this, a red navigation bar contains the breadcrumb 'arXiv.org > cs > cs.IR', a search bar with the placeholder 'Search or Article-id', and links for '(Help | Advanced search)'. A dropdown menu shows 'All papers' and a 'Go!' button. The main content area is titled 'Information Retrieval' and 'Authors and titles for recent submissions'. It lists submission dates from Friday, September 5, 2014, back to Monday, September 1, 2014. A link indicates there are 13 entries in total, with the first 25 shown. The first entry is for the paper 'On the Accuracy of Hyper-local Geotagging of Social Media Content' by David Flatow, Mor Naaman, Ke Eddie Xie, Yana Volkovich, and Yaron Kanza. It includes links to the paper's PDF and other versions, mentions 10 comments, and lists subjects: Information Retrieval (cs.IR) and Social and Information Networks (cs.SI).

Cornell University Library

We gratefully acknowledge support from the Simons Foundation and member institutions

arXiv.org > cs > cs.IR

Search or Article-id [\(Help | Advanced search\)](#)

All papers Go!

Information Retrieval

Authors and titles for recent submissions

- Fri, 5 Sep 2014
- Thu, 4 Sep 2014
- Wed, 3 Sep 2014
- Tue, 2 Sep 2014
- Mon, 1 Sep 2014

[total of 13 entries: 1-13]
[showing up to 25 entries per page: [fewer](#) | [more](#)]

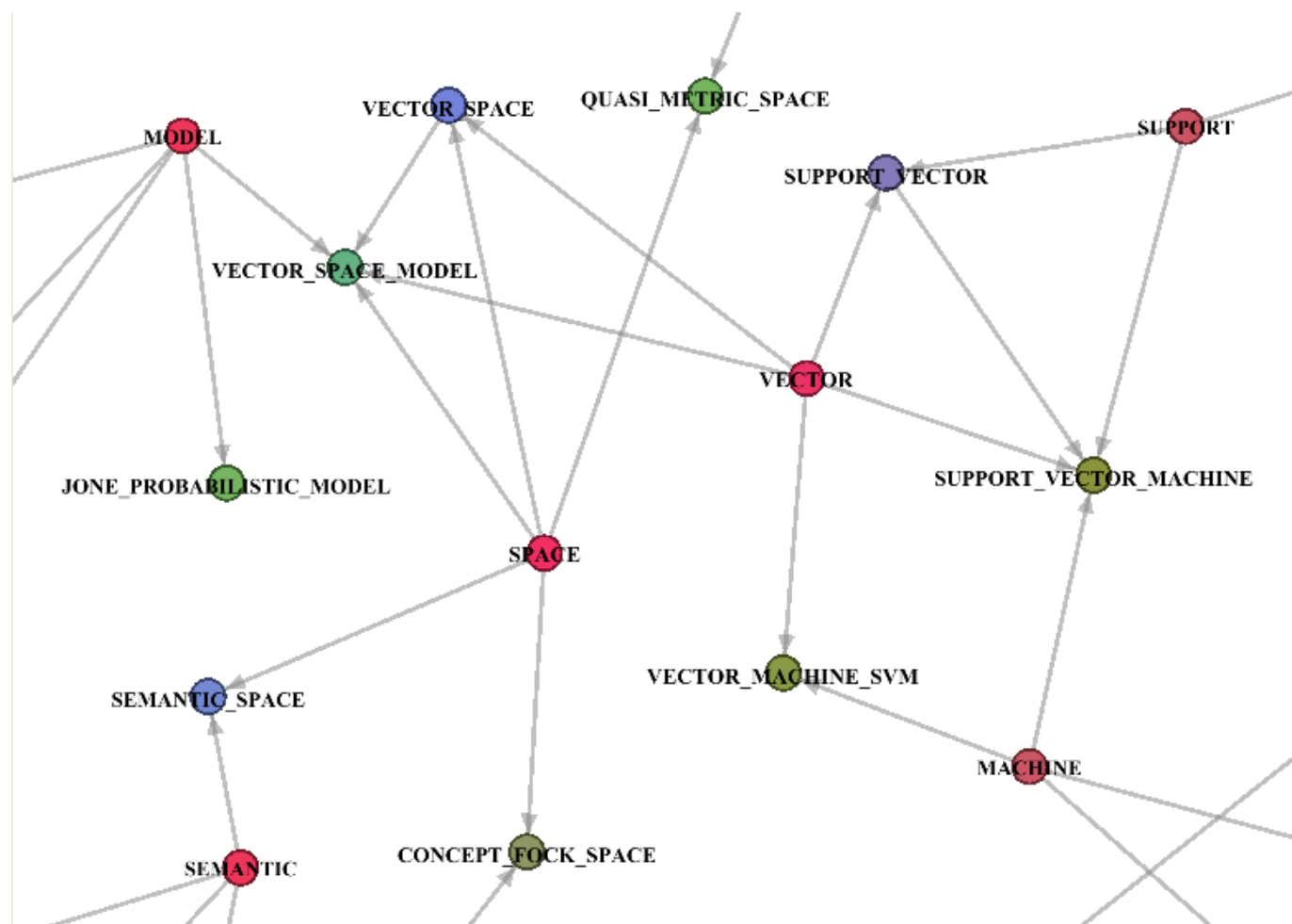
Fri, 5 Sep 2014

[1] [arXiv:1409.1461](#) [[pdf](#), [other](#)]

On the Accuracy of Hyper-local Geotagging of Social Media Content
[David Flatow](#), [Mor Naaman](#), [Ke Eddie Xie](#), [Yana Volkovich](#), [Yaron Kanza](#)
Comments: 10 pages
Subjects: **Information Retrieval (cs.IR)**; **Social and Information Networks (cs.SI)**



Візуалізація фрагменту МПІТ





Ранжирування вузлів МПІТ

За допомогою алгоритму HITS забезпечується вибір з мережі кращих «авторів» (вузлів, на які ведуть посилання) і «посередників» (вузлів, від яких йдуть посилання включення). Термін є добрим посередником, якщо від нього йдуть зв'язки на важливі словосполучення, і навпаки, добрим автором, якщо на нього ведуть зв'язки від важливих авторів. У відповідності з алгоритмом HITS для кожного вузла мережі v_j рекурсивно обчислюється його значимість як автора $a(v_j)$ і посередника $h(v_j)$ за формулами:

$$a(v_j) = \sum_i h(v_i); \quad h(v_j) = \sum_i a(v_i).$$



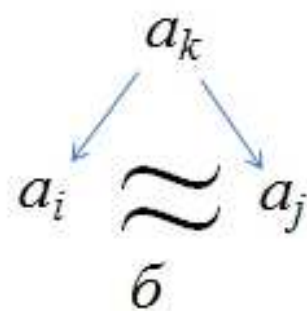
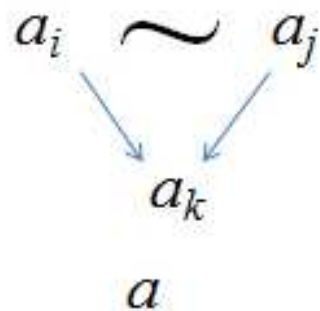
Терміни з найкращим авторством/посередництвом

Автори	Посередники
WEB_SEARCH_ENGINE	SEARCH
SEARCH_ENGINE_RESULT	ENGINE
BASED_SEARCH_ENGINE	SEARCH_ENGINE
VERTICAL_SEARCH_ENGINE	WEB
SEARCH_ENGINE_COMPANIE	RESULT
GEOGRAPHIC_SEARCH_ENGINE	SEARCH_RESULT
ENTITY_SEARCH_ENGINE	WEB_SEARCH
SEARCH_RESULT_CLUSTERING	IMAGE
TAG_SEARCH_RESUL	PAGE
SEARCH_RESULT_PRESENTATION	CLUSTERING
PERSONALIZE_SEARCH_RESULT	TAG



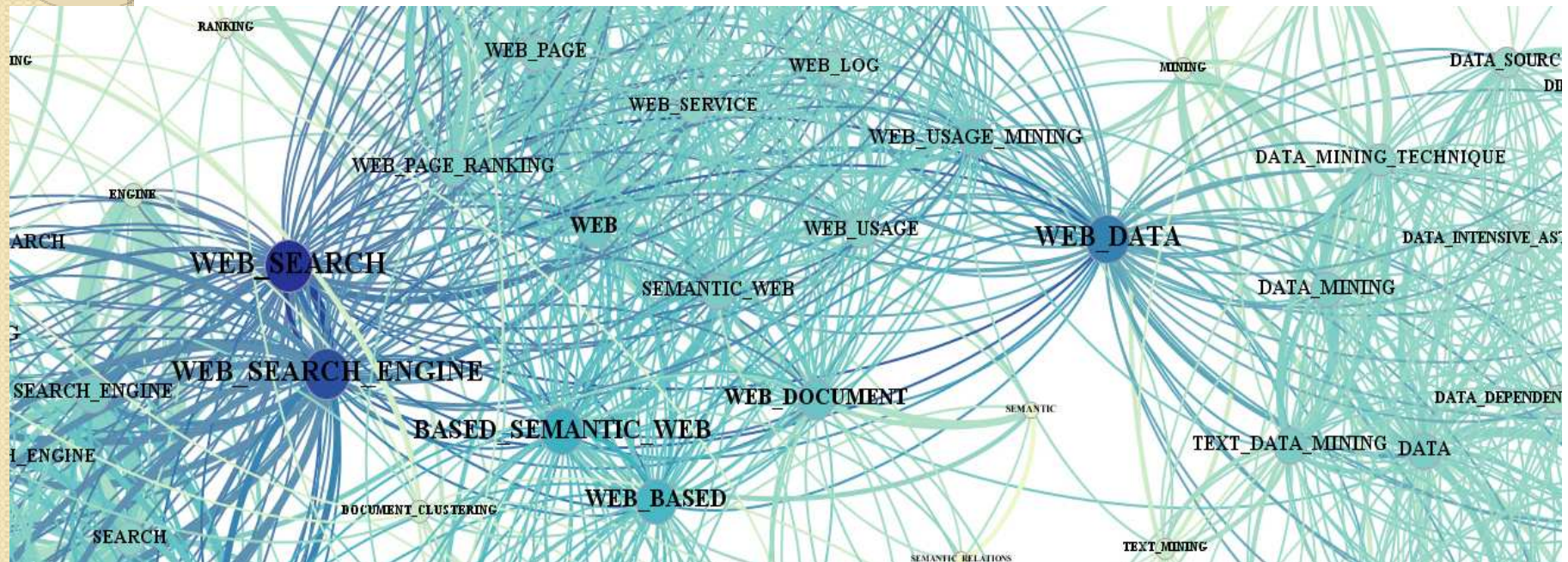
Асоціативні зв'язки

Мережу МПІТ СЕИТ можна розглядати як основу для формування зв'язків між її вузлами. Якщо позначити матрицю інцидентності МПІТ літерою A , то матриці AA^T і A^TA будуть відображувати зв'язки таких типів:





Фрагмент МПІТ з асоціативними зв'язками





Деякі висновки

- Зроблено огляд існуючих алгоритмів виявлення ключових слів.
- Запропоновано алгоритм побудови мереж природних ієрархій термінів.
- Запропоновано методи візуалізації фрагментів МПІТ.
- Запропоновано алгоритм побудови асоціативних зв'язків між термінами в МПІТ.
- Запропоновано використання алгоритму HITS для вибору найбільш важливих елементів в МПІТ.
- Мережу мови, що побудовано за допомогою запропонованої методики, можна використовувати як базу для побудови онтологій предметної області.



Міжнародна науково-технічна конференція

**ІНТЕЛЕКТУАЛЬНІ ТЕХНОЛОГІЇ ЛІНГВІСТИЧНОГО
АНАЛІЗУ**

23-24 жовтня 2013 року

Дякую за увагу!

Д.В. ЛАНДЕ

dwlande@gmail.com