

Кратність повітрообміну – це показник, який показує, скільки разів протягом години змінюється повітря в приміщенні. Враховуючи виділення діоксиду вуглецю людиною в спокої, вчені підраховали, що мінімальний об'єм вентиляції на одну людину в житлових приміщеннях повинен бути не меншим 30 м^3 за 1 годину.

Отже, кислоти та луги можуть мати негативний вплив при неправильному поводженні з ними. Дотримання правил техніки безпеки при їх використанні – гарантія безпеки власного життя. Здоров'я – це капітал, даний нам не тільки природою від народження, але і тими умовами, у яких ми працюємо.

Список використаних джерел:

1. Закон України «Про охорону праці», стаття 28.
2. Технология важнейших отраслей промышленности: Учеб. для экономич. спец. Вузов / А.М. Гинберг, Б.А. Хохлов, И.П. Дрякина и др.; Под ред. А.М. Гинберга, Б.А. Хохлова. – М.: Высш. шк., 1985. – 496 с.
3. Безопасность жизнедеятельности: Учебник / Под ред. С. В. Белова – М.: Высшая школа, 2002. – 476 с.

Кучмійов І.В.

студент,

Науковий керівник: Вовк Р.Б.

кандидат технічних наук, доцент,

Івано-Франківський національний технічний університет нафти і газу

ПЕРСПЕКТИВИ ЗАСТОСУВАННЯ ТЕХНОЛОГІЇ SEMANTIC WEB ДЛЯ ПОШУКУ ІНФОРМАЦІЇ

Пошукові системи, які виконують пошук по ключовим словам, забезпечують доступ до великої кількості сторінок для багатьох користувачів одночасно. У зв'язку з постійно зростаючою кількістю веб-сайтів зростає потреба в ретельному аналізі контенту Інтернет-документів для того, щоб звести можливість отримання нерелевантних результатів до мінімуму. Технології семантичної павутини надають можливості для вирішення проблеми релевантного пошуку. Стандартний Інтернет пошук за ключовими словами базується на пошукових технологіях, основою яких є виявлення строкової (лексичної) відповідності запитуваних термінів термінам, які присутні в Інтернет-документах.

Семантична павутина (Semantic Web) є розширенням традиційного Інтернету та націлена на спрощення пошуку і представлення інформації у вигляді, придатному для обробки комп'ютером. Дана технологія ґрунтується на елементах, побудованих з використанням стандартних мов онтологій, таких як: OWL (Web Ontology Language – мова онтології для інтернету на основі XML Web стандарту), KIF (Knowledge Interchange Format – формат обміну знаннями,

який базується на S-виразах) та CysL (онтологічна мова, яка працює на основі предикатів високого порядку) [1]. Звичайні пошукові системи працюють по принципу пошуку ключових термінів запиту в документі і не можуть використовувати його смислове значення для отримання результату, тому необхідно використовувати семантичні пошукові технології, такі як семантичні системи маршрутів (запитів) та системи типу «питання-відповідь», що використовують онтології для зберігання баз знань [2].

Документ семантичної павутини SWD (Semantic Web Document) можна розглядати як набір даних, контентом якого є або онтологія або звичайний документ розмічений за допомогою тегів. Такі Інтернет-документи можуть бути поділені на множину різних категорій, що відносяться до різних типів онтологій, які використовуються для розмітки документа. В даний час Інтернет-пошук здійснюється по ключових словах і застосовується для пошуку неструктурованих Інтернет-документів без семантичної розмітки. Найбільш популярним видом Інтернет-пошуку є булевий, який працює на виявленні комбінацій ключових слів, розділених операторами AND, OR або NOT.

Технологія ключового пошуку може бути основою для отримання SWD-документів шляхом співставлення слів-понять, які відповідають онтологічним елементам в них. Така технологія застосовується в семантичній пошуковій системі Swoogle. Відмітною її рисою є те, що вона не використовує семантику в алгоритмі виявлення, а ґрунтується на лексичних методах. Для алгоритму семантичного пошуку, лексичний і синтаксичний аналіз подібності термінів не є суттєвим, оскільки важливою є смислова схожість. Для семантичної відповідності необхідно, щоб значення запиту і документа були відомими. Якщо запит визначений формально то семантика кожного терміна може бути явно визначена. Таким чином, якщо запит представлений як онтологія, то значення кожного терміна як онтологічного поняття розкривається через використання семантичних відносин між цим поняттям та іншими онтологічними термінами. З іншого боку, якщо запит визначений неформально, наприклад на природній мові користувача, то семантика кожного терміна в запиті повинна бути розкритою. Тому залишається відкритим питання, яким чином машина зможе зрозуміти смислове значення яке малося на увазі в запиті, щоб отримати документ, найбільш близький за змістом тобто більш релевантний для користувача. В даний момент більшість пошукових систем працюють на основі підбору ключових слів, але вже існують і системи результат пошуку в яких формується на основі семантики слів, тобто пошуковий агент таких систем намагається зрозуміти семантичний сенс пошукового запиту. Частково по такому алгоритму працює система Google, яка одночасно використовує ранжування на основі ключових слів з обліком контенту гіперпосилань, що дає змогу зрозуміти сенс запиту на основі інформації, яка його оточує. В цьому випадку запити будуть зв'язуватися з фоновими знаннями зі спеціальної бази, яка в майбутньому може використовуватись як база нових алгоритмів ранжування.

Google використовує синоніми, для надання користувачам більше відповідей на пошуковий запит, з метою розширити його і дозволити підібрати

найбільш релевантний варіант. Проте застосування концепції на основі синонімів не є ідеальним рішенням, оскільки два слова в одному випадку можуть бути синонімами, а в іншому ні, що залежить від контексту, в якому вони використовуються. Наприклад, слова «coche» і «automóvil» з іспанської перекладаються як «автомобіль», а слово «carro» в латиноамериканському варіанті іспанської мови означає тільки «автомобіль», а в іспанському – «вагон» [3]. Таким чином, для підбору більш підходящих відповідей необхідно використовувати семантичний принцип аналізу слів (тобто за змістом), який дозволить краще і швидше зрозуміти пошукові запити. На основі такого принципу працює алгоритм Колібрі, в якому слова розглядаються не як окремі «елементи», а означають певні поняття, які об'єднують слова в концепції, тобто у будь-якого об'єкта може бути зв'язок з іншими об'єктами, який може змінитися в залежності від контексту. Механізми, які використовує компанія Google для визначення області пошуку, особливо важливі в процесі усунення неправильних тлумачень різних слів. З цією метою здійснюється уточнення пошуку, що підвищує точність пошукових результатів. Дана техніка є схожою до семантичного пошуку і використовується в технології «Граф знань». Об'єднання всіх цих елементів на теоретичному рівні дозволяє:

- краще зрозуміти мету пошукового запиту;
- розширити область пошуку, яка має відповідати запиту;
- спростувати видачу інформації, тобто якщо три запити по суті означають одне і те ж, то система запропонує тільки один результат;
- видавати більш точні результати пошуку завдяки семантичній обробці пошукового запиту [3].

Хоча семантична павутина сприяє пошуку релевантної інформації в мережі, але також існує і декілька невирішених проблем. Перша з них – це величезна кількість неструктурованих Інтернет-документів, які повинні бути семантично розмічені для використання семантичними пошуковими системами. Це непросте завдання, оскільки воно вимагає розвитку проблемно-орієнтованих онтологій. Другою невирішеною задачею є повністю автоматизований процес розмітки існуючих даних. З іншого боку, ефективний пошук Інтернет-документів вимагає створення пошукових запитів на основі формальних мов, а тому звичайні користувачі Інтернету повинні володіти цією мовою для створення такого роду запитів. Методи, що дозволяють автоматизувати процес перетворення запитів з природної мови користувача до формальної на даний час є мало розвинутими. Отже, семантичний пошук – це технологія майбутнього, яка являє собою багатофакторне явище і дозволяє здійснювати обробку складних та предметно-залежних запитів, які часто зустрічаються в мережі. Таким чином, в даному дослідженні розглянуто проблеми семантичного пошуку і представлено спосіб впровадження цієї технології. Також описані проблеми, що виникають в ході реалізації семантичних пошукових систем, однією з яких є проблема автоматичного перетворення запитів вільної форми до формального вигляду.

Список використаних джерел:

1. Онтология (информатика) [Електронний ресурс]. – Режим доступу: [https://ru.wikipedia.org/wiki/Онтология_\(информатика\)](https://ru.wikipedia.org/wiki/Онтология_(информатика)) Перевірено: 14.10.15.
2. Allemang D., Hendler J. Semantic Web for the Working Ontologist: Effective Modeling in RDFS and OWL// Morgan Kaufmann, 2008.
3. Gianluca F. Hummingbird Unleashed [Електронний ресурс]. – Режим доступу: <https://moz.com/blog/hummingbird-unleashed/>. Перевірено: 14.10.15.

Лунів О.В.

студент,

Науковий керівник: Вовк Р.Б.

кандидат технічних наук, доцент,

Івано-Франківський національний технічний університет нафти і газу

ЗАСТОСУВАННЯ НЕЙРОННИХ МЕРЕЖ ДЛЯ РОЗПІЗНАВАННЯ ОБРАЗІВ

На сучасному етапі розвитку інформаційні технології отримують досить широке впровадження в різні сфери життя людини, оскільки розробляються дієві підходи до вирішення різноманітних проблем і задач, що донедавна не реалізовувалися через відсутність ефективного технічного та програмного забезпечення. Зокрема, на обчислювальні структури нового типу – нейронні мережі покладається виконання інтелектуальних функцій, тобто таких, які здійснюють прийняття потрібних рішень замість людини: прогнозування банкрутств, динамічності зміни ринку, оцінки ризику страхування в економіці, діагностика кардіологічних та онкологічних захворювань в медицині, кодування і декодування інформації для безпеки різних персональних та корпоративних систем, пошук інформації, інтелектуальний аналіз даних, здійснення криміналістичної експертизи, а також ідентифікація відбитків пальців в охоронних системах [1]. Для вирішення перелічених вище задач нейронні мережі, як правило, використовують розпізнавання образів. Це технологія виділення елементів, що належать конкретному класу за допомогою виокремлення істотних ознак, із множини розмитих елементів, що відносяться до кількох класів [2], тому вивчення цього процесу і поглиблення знань про нього є досить актуальним.

Типовими задачами для розпізнавання образів є: розпізнавання цифр, букв, штрих-кодів, телефонних та автомобільних номерів, розпізнавання обличчя, зображень, мови та розпізнавання складних технологічних, інфраструктурних, економічних явищ і процесів. Розпізнавання образів здійснюється декількома методами, зокрема:

- методом порівняння основних характеристик образу з певним еталоном для класу образів;