



ҚАЗАҚСТАН РЕСПУБЛИКАСЫ
БІЛІМ ЖӘНЕ ҒЫЛЫМ МИНИСТРЛІГІ
МИНИСТЕРСТВО ОБРАЗОВАНИЯ И НАУКИ
РЕСПУБЛИКИ КАЗАХСТАН
MINISTRY OF EDUCATION AND SCIENCE
OF THE REPUBLIC OF KAZAKHSTAN



Л. Н. ГУМИЛЕВ АТЫНДАҒЫ
ЕУРАЗИЯ ҰЛТТЫҚ УНИВЕРСИТЕТІ
ЕВРАЗИЙСКИЙ НАЦИОНАЛЬНЫЙ
УНИВЕРСИТЕТ ИМ. Л. Н. ГУМИЛЕВА
GUMILYOV EURASIAN
NATIONAL UNIVERSITY



Студенттер мен жас ғалымдардың
«Ғылым және білім - 2015»
атты X Халықаралық ғылыми конференциясының
БАЯНДАМАЛАР ЖИНАҒЫ

СБОРНИК МАТЕРИАЛОВ
X Международной научной конференции
студентов и молодых ученых
«Наука и образование - 2015»

PROCEEDINGS
of the X International Scientific Conference
for students and young scholars
«Science and education - 2015»

УДК 001:37.0
ББК72+74.04
Ғ 96

Ғ96

«Ғылым және білім – 2015» атты студенттер мен жас ғалымдардың X Халық. ғыл. конф. = X Межд. науч. конф. студентов и молодых ученых «Наука и образование - 2015» = The X International Scientific Conference for students and young scholars «Science and education - 2015». – Астана: <http://www.eni.kz/ru/nauka/nauka-i-obrazovanie-2015/>, 2015. – 7419 стр. қазақша, орысша, ағылшынша.

ISBN 978-9965-31-695-1

Жинаққа студенттердің, магистранттардың, докторанттардың және жас ғалымдардың жаратылыстану-техникалық және гуманитарлық ғылымдардың өзекті мәселелері бойынша баяндамалары енгізілген.

The proceedings are the papers of students, undergraduates, doctoral students and young researchers on topical issues of natural and technical sciences and humanities.

В сборник вошли доклады студентов, магистрантов, докторантов и молодых ученых по актуальным вопросам естественно-технических и гуманитарных наук.

УДК 001:37.0
ББК 72+74.04

ISBN 978-9965-31-695-1

©Л.Н. Гумилев атындағы Еуразия
ұлттық университеті, 2015

$$\begin{aligned} \theta_1 &\rightarrow \max \\ \sum_{j \neq k} (b_j - a_k) y_j &\geq \theta_1 \quad k=1, \dots, n. \\ I^T Y &= 1, Y \geq 0 \end{aligned}$$

Эту задачу можно записать так

$$\min_k \left(\sum_{j \neq k} y_j b_j - (1 - y_k) a_k \right) : I^T Y = 1, Y \geq 0 \rightarrow \max_Y.$$

То есть 1-ый игрок хочет сформировать портфель, максимизирующий свой наименьший выигрыш.

Оптимальное решение этой задачи дает гарантированный результат, т.е. независимо от развития событий, величина выигрыша заведомо не меньше $\theta_{\max}^1$.

Все выводы теории матричных игр здесь применимы: $\theta_{\min}^2 = \theta_{\max}^1$.

Список использованных источников

4. Шапкин А.С., Шапкин В.А. Экономические и финансовые риски. Оценка, управление, портфель инвестиций / М.: «Дашков и К^о», 2011.
5. Нуртазина К.Б. Оптимизация портфеля ценных бумаг и управление в условиях неопределенности. М.: ГУУ, 2011.

УДК 519.8

ИНТЕЛЛЕКТНОЕ МОДЕЛИРОВАНИЕ ДИНАМИЧЕСКИХ ПРОЦЕССОВ В УСЛОВИЯХ НЕОПРЕДЕЛЕННОСТИ

Сейтенов Алтынбек Серикович

altynbekss@gmail.com

Магистрант 2 курса специальности математического и компьютерного моделирования
ЕНУ им.Л.Н.Гумилева, г.Астана, Казахстан
Научный руководитель Танирбергенов А.Ж.

АННОТАЦИЯ: В статье рассмотрены подходы и механизмы содействия нечетких запросов с реляционными базами данными, базирующиеся на теории нечетких множеств Лотфи Заде. Нечеткая система, которая поддерживает вычисления и обработки данных будет основываться на основе интегрированных нечетких утилит в Microsoft .NET и Ms Sql Server. Так же рассмотрено формирование оценок с помощью лингвистических переменных.

КЛЮЧЕВЫЕ СЛОВА: нечеткие запросы, лингвистические переменные, реляционные базы данных, SQL запросы.

В наше время по мере увеличения объема данных, интеллектуальная обработка и анализ данных стало более самым интересным, полезным и популярным, так и при использовании в частных бизнес-системах, так и в открытой сети Интернет. Так как некоторые форматы информационной неопределенности, традиционными технологиями обработать данные затруднительно, и в таких ситуациях могут быть полезны теория нечетких множеств [1].

Основная часть базы данных были разработаны для хранения данных, и конечно с важной способностью получать доступ к данным после их сохранения. Есть многие разновидности систем управления базами данных. Самым распространенным является

реляционные модели, с использованием структурированного языка запросов SQL. Бывают документ-ориентированные модели, например: MongoDB и Apache CouchDB и, конечно же графовые модели, такие как RDF. Все они имеют возможность хранить детально и точно данные, а также запрашивать данные, хранящиеся в них.

Запрос данных в строгой форме и формате позволяет получить только определенный и конкретную информацию. Более высоким уровнем, чем простой запрос, является статистический анализ, основанный на определенных подсчетах (процентных отношениях и пропорциях). С помощью статистики вы можете предсказывать возможные варианты будущего, но для этого используются, как правило, закрытые данные, которые хранятся в вашей базе данных.

На данный момент существует ряд методов интеллектуального анализа данных, многие из которых специально предназначены для статистического анализа и извлечения возможной новой информации из существующих данных. В связи с актуальностью нечеткой математики, растет количество алгоритмов и методик анализа данных способных обрабатывать неопределенность и расплывчатость данных, и эти предложенные системы анализов позволяют выдавать неоднозначные и расплывчатые результаты для структурированных запросов.

В реальной жизни нас повсюду и постоянно окружают неопределённости и неточности. Например, с определением «У него хорошая зарплата». Это и есть неопределенность, для каждого понятие «хорошая зарплата» разное и ликвидно, а для машины и ЭВМ этот термин совсем незнаком. Ваши коллеги могут сказать, что он получает «около 150 000тг в месяц». Слово «около» говорит о предполагаемой неопределенности, поскольку зарплата указано неточно [2].

И первый кто начал рассматривать этот вопрос, отец-основатель целой научной теории, фундаментальных трудов «Fuzzy Sets» Лотфи Заде. Но при всех заслугах Л. Заде, важный вклад внесли и его последователи. Например, английский математик Э. Мамдани. Он разработал алгоритм, который был предложен в качестве метода для управления паровым двигателем. Этот алгоритм, основанный на нечетком логическом выводе, позволил избежать чрезмерного большого объема вычислений, и был по достоинству оценен специалистами. Этот алгоритм в настоящее время получил наибольшее практическое применение в задачах нечеткого моделирования.

Чтобы помочь преодолеть проблемы, связанные с неопределённостью данных, специализированные производители предлагают разнообразные решения, чтобы помочь выявить и исправить тонкие различия в идентичных данных с помощью процессов, основанные на алгоритме нечетких множеств. Одна из этих модулей была интегрирована в MS SQL Server 2008, услуг SQL Server Integration (SSIS).

Для создания и хранения данных, с которыми будут обрабатываться пример – Формула 1, мы использовали среду MS SQL Server 2008.

```

CREATE TABLE fuzzyLookupSource
( firstName varchar(10),
  LastName varchar(10),
  BirthDate datetime )
INSERT INTO fuzzyLookupSource
VALUES(' Altynbek', 'Seitenov', 1992 - 02 - 27)
INSERT INTO fuzzyLookupSource
VALUES(' Aigul', 'Ibadildina', 1964 - 05 - 2'5)
INSERT INTO fuzzyLookupSource
VALUES(' Serik',    'Seitenov', 1963 - 10 - 24 )
INSERT INTO fuzzyLookupSource
VALUES(' Kuanysh', 'Omarov', 1991 - 09 - 15)
INSERT INTO fuzzyLookupSource
VALUES(' Altynbek', 'Kunanbay', 1992 - 12 - 23)

```

Формула 1- создание базы данных

После того как данные были созданы, с использованием SQL Server 2008 Business Intelligence Studio, создается новая служба SQL Server Integration Services (SSIS). Далее наши созданные данные маршрутизируем в интегрированный утилит «Fuzzy Lookup Transformation», для определения «нечетких совпадений».

Так же устанавливаются пороговые критерии непосредственно в Fuzzy Lookup и выбираются столбцы с условие на базе методов _Similarity и _Confidence. Условное разбиение вычисляет, и разбивают строки в наборе данных. Рисунок 1

Поряд...	Имя выхода	Условие
1	Solid Matched	_Similarity > 0.85 && _Confidence > 0.8
2	Likely Matched	_Similarity > .65 && _Confidence > 0.7

Имя выхода по умолчанию:

Рисунок 1. Условное разбиение

Далее данные разделяются по классам совпадений и записываются с производные столбцы. Производные столбцы обновляют столбцы с помощью выражениями, которые мы их связали. И так мы проводим слияние наборов данных и загружаем виде нового запроса в реляционные базы данных с помощью поставщика OLE DB. Рисунок 2.

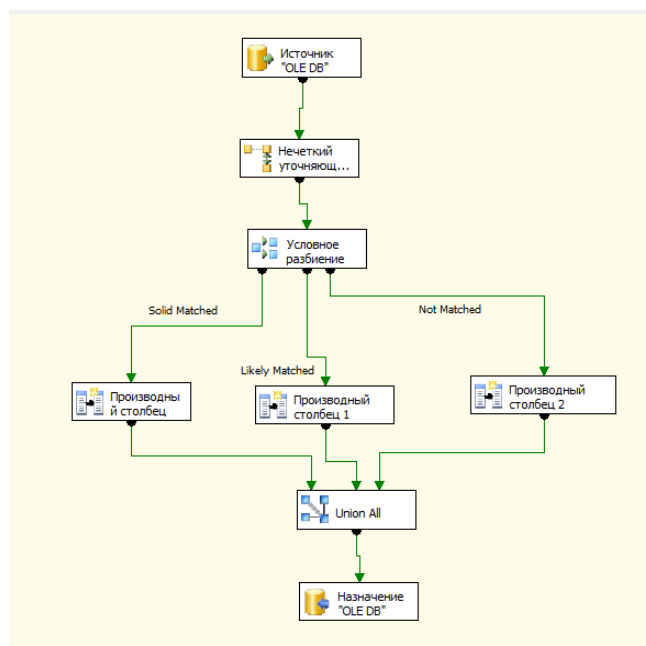


Рисунок 2. Поток управления данными

Для получения лучшего представления использовалась форма (изображена на рисунке 3), написанная на языке программирования C# в интегрированной среде разработки Microsoft Visual Studio 2010. Она дает возможность хороший способ обзора данных, полученные в виде результата и экономит много времени для сравнения.

Рисунок 3. Форма сравнения

В целом, нечеткая классификация данных с помощью служб интеграции SQL Server выглядит очень гибкой, и простой в управлении. Функции, которые создают последовательные и воспроизводимые отчеты организаций, могут использоваться и легко удовлетворять собственные нужды. Превосходная смесь функциональности, а также привлекательности для оценки данных.

Список использованных источников

1. Круглов В.В., Дли М.И., Голунов Р.Ю. Нечеткая логика и искусственные нейронные сети. Москва 2001, С. 55-103
2. Леоненков А.В. Нечеткое моделирование в среде MATLAB и fuzzyTECH. БХВ-Петербург, 2003, С. 43-69
3. Стилмен Э., Грин Дж. Изучаем C#. Москва 2012, С. 39-105