



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

ПОДСЕКЦИЯ 2.4 - ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

УДК 004.9

ПРЕИМУЩЕСТВО ИСПОЛЬЗОВАНИЯ ПРОГРАММЫ-ПЕРЕВОДЧИКА PRAGMA

Базарова Алтынай Кемелбеккызы

Магистрант кафедры «Информатики и информационной безопасности»

ЕНУ им. Л.Н.Гумилева, Астана, Казахстан

Научный руководитель – Б. Андасова

В ее основе работы программы переводчика лежит алгоритм перевода – последовательность однозначно и строго определенных действий над текстом для нахождения соответствий в данной паре языков L1 – L2 при заданном направлении перевода (с одного конкретного языка на другой). Обычные словари и грамматики разных языков не применимы для машинного перевода, так как описывают значения слов и грамматические закономерности в нестрогой форме, никак не приемлемой для «машинного» использования. Следовательно, нужна формальная грамматика языка, т.е. логически непротиворечивая и явно выраженная (безо всяких подразумеваний и недомолвок). Как только начали появляться формальные описания различных областей языка – прежде всего морфологии и синтаксиса, – наметился прогресс и в разработке систем автоматического перевода. Чтобы успешно работать, система машинного перевода включает в себя, во-первых, двуязычные словари, снабженные необходимой информацией (морфологической, относящейся к формам слова, синтаксической, описывающей способы сочетания слов в предложении, и семантической, т.е. отвечающей за смысл), а во-вторых – средства грамматического анализа, в основе которых лежит какая-нибудь из формальных, т.е. строгих, грамматик. Наиболее распространенной является следующая последовательность формальных операций, обеспечивающих анализ и синтез в системе машинного перевода.

1. На первом этапе осуществляется ввод текста и поиск входных словоформ (слов в конкретной грамматической форме, например дательного падежа множественного числа) во входном словаре (словаре языка, с которого производится перевод) с сопутствующим морфологическим анализом, в ходе которого устанавливается принадлежность данной словоформы к определенной лексеме (слову как единице словаря). В процессе анализа из формы слова могут быть получены также сведения, относящиеся к другим уровням организации языковой системы, например, каким членом предложения может быть данное слово. В школьном грамматическом разборе предложения мы опираемся и на значения слов, составляющих предложение (например, отыскивая подлежащее, задаем вопрос: о чем говорится в предложении?). Для машины же совмещение двух этих операций – и грамматического разбора, и обращения к смыслу слов – задача трудная. Лучше сделать синтаксический анализ не зависящим от смысла слов, а словарь использовать на других этапах перевода.

Что такое независимый синтаксический анализ, можно понять, если попытаться разобрать фразу, из которой «убраны» значения конкретных слов. Блестящим образцом фразы такого рода является придуманное академиком Л. В. Щербой предложение: Глокая куздра штетко будланула бокра и кудрячит бокрѣнка. Бессмысленная фраза? Как будто да: в русском языке нет слов, из которых она состоит (кроме союза и). И все же в какой-то степени мы ее понимаем: «куздра» – это существительное (мы даже можем предположить, что оно обозначает какое-то животное), «глокая» – определение к нему, «будланула» – глагол-сказуемое (похожий на толкнула, боднула), «штетко» – скорее всего, обстоятельство образа действия (что-то вроде сильно, резко), «бокра» – это прямое дополнение («будланула» кого? – «бокра») и т. д.

То есть машина осуществляет синтаксический анализ предложения без опоры на значения составляющих его слов, с использованием информации только об их грамматических свойствах. В результате синтаксического анализа возникает синтаксическая структура, которая изображается в виде дерева зависимостей: «корень» – сказуемое, а «ветви» – синтаксические отношения его с зависимыми словами. Каждое слово предложения записывается в своей словарной форме, а при ней указываются те грамматические характеристики, которыми обладает это слово в анализируемом предложении.

2. Следующий этап включает в себя перевод идиоматических словосочетаний, фразеологических единств или штампов данной предметной области (например, при англо-русском переводе обороты типа *in case of*, *in accordance with* получают единый цифровой эквивалент и исключаются из дальнейшего грамматического анализа); определение основных грамматических (морфологических, синтаксических, семантических и лексических) характеристик элементов входного текста (например, числа существительных, времени глагола, их роли в данном предложении и пр.), производимое в рамках входного языка; разрешение неоднозначности (скажем, англ. *round* может быть существительным, прилагательным, наречием, глаголом или же предлогом); анализ и перевод слов. Обычно на этом этапе однозначные слова отделяются от многозначных (имеющих более одного переводного эквивалента в выходном языке), после чего однозначные слова переводятся по спискам эквивалентов, а для перевода многозначных слов используются так называемые контекстологические словари, словарные статьи которых представляют собой алгоритмы запроса к контексту на наличие/отсутствие контекстных определителей значения.

3. Окончательный грамматический анализ, в ходе которого доопределяется необходимая грамматическая информация с учетом данных выходного языка (например, при русских существительных типа *сани*, *ножницы* глагол должен стоять в форме множественного числа, притом, что в оригинале может быть и единственное число).

4. Синтез выходных словоформ и предложения в целом на выходном языке. Здесь не получится обойтись простым переводом «узлов» дерева на другой язык. Синтаксис каждого языка устроен на свой лад: то, что в русском предложении – подлежащее, в другом языке может (или должно) быть выражено дополнением, а дополнение, наоборот, должно преобразоваться в подлежащее; то, что в одном языке обозначается группой слов, переводится на другой всего одним словом и т. д. Так, при переводе русской фразы «У меня была интересная книга» на английский язык глагол «быть» надо перевести глаголом *to have* – «иметь», сочетание «у меня» преобразовать в подлежащее *I* («я»), а слово «книга», которое в русском языке – подлежащее, по-английски должно стать прямым дополнением: *I had an interesting book* (буквально: «Я имел интересную книгу»). В связи с этим в машинную память помимо наборов синтаксических правил для каждого языка «вкладывают» и правила преобразования синтаксических структур. К этому добавляют правила перехода от уже преобразованной структуры к предложению того языка, на который делается перевод. Такой переход от структуры к реальному предложению называется синтаксическим синтезом.

В зависимости от особенностей морфологии, синтаксиса и семантики конкретной языковой пары, а также направления перевода общий алгоритм перевода может включать и другие этапы, а также модификации названных этапов или порядка их следования, но вариации такого рода в современных системах, как правило, незначительны. Анализ и синтез могут производиться как пофразно, так и для всего текста, введенного в память компьютера; в последнем случае алгоритм перевода предусматривает определение так называемых анафорических связей (такова, например, связь местоимения с замещаемым им существительным – скажем, местоимения *им* со словом местоимения в самом этом пояснении в скобках).

Для решения проблемы многозначности слова используется анализ контекста. Дело в том, что каждое из нескольких значений многозначного слова в большинстве случаев реализуются в своем наборе контекстов. То есть у каждого из «конкурирующих» (при интерпретации) значений – свой набор контекстов. И именно вот эта зависимость значения

от окружения позволяет слушающему понять высказывание правильно. Для правильного понимания высказывания необходимо в полной мере учитывать также правила обусловленности выбранного значения лексическим окружением (действующие при «фразеологической» интерпретации слова), правила обусловленности выбранного значения семантическим контекстом (так называемые законы семантического согласования) и правила обусловленности выбранного значения грамматическим (морфолого-синтаксическим) контекстом. То есть для решения проблемы «моносемизации» слов при автоматическом переводе основой служит изучение и тщательное описание закономерностей лексической, семантической и грамматической сочетаемости. При этом правила такой сочетаемости достаточно подробно описываются в словарях – а именно, (а) с мощным охватом лексики, но весьма бегло и нетщательно, а также весьма имплицитно это делается в традиционной лексикографии; и, с другой стороны, (б) в выборочном порядке (со слабым охватом лексики), но зато весьма аккуратно и тщательно, и довольно-таки эксплицитно это делается в работах по «толково-комбинаторной» лексикографии (последних сорока лет).

Действующие системы машинного перевода, как правило, ориентированы на конкретные пары языков (например, французский и русский или японский и английский) и используют, как правило, переводные соответствия либо на поверхностном уровне, либо на некотором промежуточном уровне между входным и выходным языком. Качество машинного перевода зависит от объема словаря, объема информации, приписываемой лексическим единицам, от тщательности составления и проверки работы алгоритмов анализа и синтеза, от эффективности программного обеспечения. Современные аппаратные и программные средства допускают использование словарей большого объема, содержащих подробную грамматическую информацию. Информация может быть представлена как в декларативной (описательной), так и в процедурной (учитывающей потребности алгоритма) форме.

В практике переводческой деятельности и в информационной технологии различаются два основных подхода к машинному переводу. С одной стороны, результаты машинного перевода могут быть использованы для поверхностного ознакомления с содержанием документа на незнакомом языке. В этом случае он может использоваться как сигнальная информация и не требует тщательного редактирования. Другой подход предполагает использование машинного перевода вместо обычного «человеческого». Это предполагает тщательное редактирование и настройку системы перевода на определенную предметную область. Здесь играют роль полнота словаря, ориентированность его на содержание и набор языковых средств переводимых текстов, эффективность способов разрешения лексической многозначности, результативность работы алгоритмов извлечения грамматической информации, нахождения переводных соответствий и алгоритмов синтеза. На практике перевод такого типа становится экономически выгодным, если объем переводимых текстов достаточно велик (не менее нескольких десятков тысяч страниц в год), если тексты достаточно однородны, словари системы полны и допускают дальнейшее расширение, а программное обеспечение удобно для постредактирования.

За рубежом эксплуатируется целый ряд систем машинного перевода. Наиболее известной из их числа является система Systran, разработанная и поддерживаемая компанией Systran Software Inc, используемая службой машинного перевода при комиссии Европейского союза.

Появилась возможность воспользоваться услугами автоматических переводчиков непосредственно в Сети: www.alphaworks.ibm.com/aw.nsf/html/mt; www.freetranslation.com; www.transtlate.ru; www.logomedia.net/text.asp; www.foreignword.com/Tools/transnow.htm; babelfish.altavista.com/translate.dyn; infinet.reverso.net/traduire.asp; www.t-mail.com.

С начала 1990-х годов на рынок систем ПК выходят отечественные разработки.

В июле 1990 года на выставке PC Forum в Москве была представлена первая в России коммерческая система машинного перевода под названием PROMT (PROgrammer's Machine Translation). В 1991 г. было создано ЗАО «ПРОект МТ», и уже в 1992 г. компания «ПРОМТ»

выиграла конкурс NASA на поставку систем МП (ПРОМТ была единственной неамериканской фирмой на этом конкурсе).

Несмотря на такую долгую историю, фактически всеми системами осуществляется перевод только на уровне поверхностного синтаксиса, поскольку еще не разработаны (по всей видимости) эффективные модели формального представления смысла, носителем которого должен выступать язык-посредник – интерлингва, хотя для отдельных узких отраслей такие модели строятся (например, METEO и LingoWare). Специалисты связывают построение адекватных систем МП с развитием искусственного интеллекта: машина сможет переводить с одного языка на другой, когда научится думать, как человек.

Другой путь совершенствования МП, более доступный на современном этапе, – составить корпус соответствий на двух языках. Можно предположить, что такие работы ведутся, и многими разными командами, но их действия не скоординированы, и потому результат слишком мал.

Критики современных систем МП полагают, что установка на жанровую ограниченность (научить машину сначала понимать совсем простые, специально отобранные тексты) на практике привела к тому, что задача моделирования естественного языка фактически уступила место задаче моделирования ограниченных (и крайне примитивных) подязыков отдельных отраслей знания. При этом наилучшего результата на этом пути, как известно, достигла канадская система TAUM-METEO, отлично выполняющая задачу англо-французского перевода сводок погоды. Простейшим видом систем такого рода являются автоматические разговорники для туристов, предлагающие пользователю более или менее разнообразные «меню» стандартных вопросов и ответов на двух или нескольких языках.

Существующий в настоящее время «словоцентрический» подход (когда машина выбирает и переводит главным образом отдельные слова) объясняется тем, что выделяется то, что легко выделить (слова разделены пробелами), и, соответственно, это переводится. Однако человек (в том числе тот, который занимается переводом) имеет дело с текстом, когда отдельное предложение приобретает смысл как часть более широкого контекста: соседние предложения определяют и объясняют многие невыраженные или неоднозначные элементы каждого отдельного высказывания. На настоящем же этапе часто самыми удобными для понимания оказываются такие системы МП, которые выполняют перевод пословно: фраза корявая, но видно, как она получилась, и, если есть поддержка в виде знания исходного языка, легко догадаться, что же было в оригинале, и увидеть, какие слова переведены неверно. Те системы, которые переводят текст пословно, зачастую оказываются удобнее: видно, откуда фраза взялась. Если хотя бы поверхностно знать язык оригинала, можно понять, что же было в первоначальном варианте, и какие слова переведены неверно. Системы МП, которые обрабатывают фразу синтаксически, избегая «корявости», часто выдают гладкие, но совершенно невразумительные переводы.

ТМ – это база данных, где хранятся выполненные переводы. Технология ТМ работает по принципу накопления: в процессе перевода в ТМ сохраняется исходный сегмент (предложение) и его перевод. При обработке нового текста, поступившего на перевод, система сравнивает каждое его предложение с сохраненными в базе сегментами. Если идентичный или подобный исходному сегмент найден, то перевод этого сегмента отображается вместе с переводом и указанием совпадения в процентах. Слова и фразы, которые отличаются от сохраненного текста, выделяются подсветкой. Таким образом, переводчику остается перевести только новые сегменты и отредактировать частично совпадающие. Каждое изменение или новый перевод сохраняются в ТМ. А в результате нет необходимости дважды переводить одно и то же предложение.

С другой стороны, при работе с крупными проектами переводчик сталкивается с проблемой согласованного применения терминологического глоссария в ходе длительного проекта или быстрого повторного использования ранее переведенного текста. По своей природе подобные рутинные задачи сравнительно легко (в отличие от машинного перевода) формализуются и программируются.

Каждая запись базы данных ТМ представляет собой единицу (предложение или абзац) параллельных текстов (как правило, на двух языках). Такая база данных хранит предыдущие переводы с целью их возможного повторного использования и решения задач быстрого поиска по содержанию. Несмотря на то что программы, оснащенные памятью перевода, называются системами автоматизированного перевода (CAT, computer-aided/assisted translation), их не следует путать с программами машинного перевода (machine translation) – память перевода ничего не переводит сама по себе, в то время как машинный перевод основан на генерации переводов по результатам грамматического разбора исходного текста.

Как правило, запись памяти перевода состоит из двух сегментов: на исходном (source) и конечном (target) языках. Если идентичный (или похожий) сегмент на исходном языке встречается в тексте, сегмент на конечном языке будет найден в памяти перевода и предложен переводчику в качестве основы для нового перевода. Автоматически найденный текст может быть задействован как есть, отредактирован или полностью отвергнут. Большинство программ используют алгоритм нечеткого соответствия (fuzzy matching), существенно улучшающий их функциональные возможности, поскольку в этом случае можно находить предложения, лишь отдаленно напоминающие искомые фразы, но, тем не менее, пригодные для последующего редактирования.

Преимущества от использования такого программного обеспечения поначалу могут быть неочевидны – однако по мере наполнения базы данных результаты автоматической подстановки основ для перевода будут становиться все более точными и регулярными.

Архитектура автоматизированной системы и ее функциональные возможности могут различаться. Средства поиска могут работать как с целыми сегментами, так и с отдельными словами или фразами, позволяя переводчику выполнять терминологический поиск. В систему также включают отдельную программу для работы с глоссарием, содержащим утвержденные для применения в проекте термины. Некоторые системы работают с программами машинного перевода. Основной рабочий интерфейс либо встраивается непосредственно в имеющийся текстовый процессор, такой как Word, либо представляет собой отдельный редактор. В состав системы обязательно включают фильтры для импорта-экспорта файлов различных форматов. Кроме того, многие системы, если не все, имеют средство для добавления в память перевода сегментов из, как правило, имеющихся у переводчика старых переведенных файлов.

То, что применимо к понятию «обучение языку», применимо и к «Translation Memories».

- «Пустая» система запоминает термины и предложения.
- Строится «память переводов» – Translation Memory (TM).
- TM становится «языковой памятью» по продукту или по деятельности компании в целом.

Системы TM: SDLX, TRADOS, Deja Vu, Star Transit, Trans Suite 2000, WordFast, WordFisher, ACROSS.

Технологии МП и ТМ друг друга дополняют, но никак не дублируют. Система МП готова к использованию сразу после установки (хотя это не исключает того, что в процессе работы пользователю захочется что-то изменить в словаре, алгоритмах перевода и т.п.). Систему ТМ необходимо специально настраивать на перевод текстов в какой-то конкретной области, и чем больше эти тексты друг на друга похожи (например, такая система используется для перевода стандартных договоров), тем меньше времени требуется для настройки.

В связи с этим совершенно логично появление гибридов – example-based machine translation – программ, объединяющих системы машинного перевода и ТМ (например, компания «ПРОМТ» создала интегрированную технологию PROMT Term и PROMT For TRADOS, которая объединяет систему ТМ TRADOS и систему машинного перевода – PROMT XT Professional). PROMT For TRADOS (P4T) предназначена для интеграции системы машинного перевода PROMT и системы ТМ TRADOS:

- перевод в системе TRADOS;

- перевод в системе PROMT не найденных в ТМ сегментов;
- вставка переведенных PROMT сегментов в ТМ.

Схема автоматизированной цепочки перевода на основе интегрированной технологии PROMT-TRADOS

Применение интегрированной технологии делает процесс перевода больших массивов документации управляемым и повышает его экономическую эффективность.

Пример реализации проектов с применением интегрированной технологии PROMT-TRADOS.

Предположим, необходимо осуществить перевод инструкции к мини-АТС.

1. На первом этапе применяется программа PROMT TerM. Документы анализируются, и выявляется основная терминология, которая заносится в словари системы машинного перевода PROMT.

2. Выполняется машинный перевод (МП) с подключенным словарем, продолжается терминологическая работа по коррекции словаря.

3. Результаты МП корректируются и заносятся в ТМ переводимого документа.

4. Таким образом пользователь получает:

- терминологический словарь;
- переведенный документ;
- соответствующую переведенному документу ТМ, которая может быть использована в дальнейшей работе с документами подобного рода.

Чтобы выбрать программу-переводчик, надо знать, насколько хорошо она работает. Для проверки PROMT и Pragma мы взяли по короткому тексту на английском и русском языках и предложили их для перевода обоим программам.

Вот как переводил текст с английского на английский PROMT и Pragma. Текст был следующий: «Linguistic research based on real-life data has become more and more prominent during the last years. As a consequence, the need for corpora that are (i) large and (ii) richly annotated has grown as well. First, corpora that consist of real-life data such as newspaper texts or recorded dialogues must be large enough to offer enough instances of the phenomena under study. Second, many linguistic phenomena involve factors of different linguistic domains; for instance, word order in German is supposed to depend, among other things, on grammatical functions (syntax), thematic roles (semantic), information structure, and intonation (phonetics). Investigations of such phenomena require corpora that are annotated with detailed information at various linguistic levels».

Как видно, текст достаточно формальный, без художественных пассажей, непреодолимых для машинного мозга на его современном этапе. Вот как это переводит Promt: «Лингвистическое исследование, основанное на реальных данных стало более видным в течение прошлых лет. Как следствие, потребность в корпусах, которые являются (i) большими и (ii) богато аннотируемым, росла также. Сначала, корпуса, которые состоят из реальных данных, типа газетных текстов и записанных диалогов, должны быть достаточно большими, чтобы предложить достаточно многие случаи явлений при исследовании.

Во-вторых, много лингвистических явлений вовлекают факторы различных лингвистических областей; например, порядок слов на немецком языке должен зависеть, между прочим, на грамматических функциях (синтаксис), тематические роли (семантическая), информационная структура, и интонация (фонетика). Исследования таких явлений требуют корпусов, которые аннотируются с подробной информацией на различных лингвистических уровнях».

Вот что думает по этому поводу Pragma: «Лингвистическое исследование, основанное на данных реальной жизни, стало еще и еще выдающимся в течение прошлых лет. Как следствие, потребность в массивах, которые (i) большие и (второй) богато аннотировано, росла также. Для начала, массивы, которые состоят из данных реальной жизни как например газетные тексты и записанные диалоги, должны быть достаточно большими, чтобы предложить достаточно образцов явлений для изучения. Во-вторых, много

лингвистических явлений вовлекают факторы различных лингвистических областей; например, в немецком языке предполагают, что заказ слова зависит, среди других вещей, на грамматических функциях (синтаксис), тематических ролях (семантический), информационной структуре, и интонации (фонетика). Исследования таких явлений требуют массивы, которые аннотированы с детальной информацией на различных лингвистических уровнях».

С русского на английский эти программы переводили следующий текст: «Германские физики приступили к созданию изобретенного Альбертом Эйнштейном абсорбционного холодильника, работающего на спирте. Мало кто знает, что на счету основателя современной физики было два десятка патентов на различные технические изобретения. Большинство из них касаются холодильных машин, которые великий теоретик создавал соавторстве с крупнейшим ядерным физиком Лео Сцилардом, одним из создателей первого ядерного реактора.

Компрессор сконструированного им в 1925 году абсорбционного холодильника, в отличие от наиболее распространенных в настоящее время компрессионных, не имел движущихся частей (компрессора). В качестве хладагента в нем использовался не фреон, а бутанили аммиак. В представленном в 1927 году на Лейпцигской ярмарке усовершенствованном холодильнике для охлаждения использовался уже метиловый спирт. Несмотря на то что этот холодильник по свидетельству очевидцев «выл, как шакал», один из электротехнических концернов купил патент на него за 12 тысяч марок. Но в серийное производство воющий холодильник так никогда и не был запущен».

Prompt предложил следующую версию: «The German physicists have started creation invented by Albert Einstein absorption refrigerator working on spirit. Very few people knows, that on the account of the founder of modern physics was two ten patents for various technical inventions. The majority of them concern refridgerators which the great theorist created in the co-authorship with the largest nuclear physicist Leo Stsilardom, to one of founders of the first nuclear reactor.

The compressor designed by them in 1925 absorption refrigerator, unlike the most widespread now компрессионных, had no moving parts (compressor). As a coolant in it freon, and butane or ammonia was used not. In presented in 1927 at the Leipzig fair the advanced refrigerator for cooling already methyl spirit was used. In spite of the fact that this refrigerator under the certificate of eyewitnesses « howled, as the jackal », one of electrotechnical concerns has bought the patent for it for 12 thousand marks. But in a batch production the howling refrigerator so never also has not been started».

Pragma тот же самый кусочек поняла так: «The German physicists began creation invented Albert by Einstein absorption refrigerator, working on an alcohol. The few knows that on the account of founder of modern physics there were two ten patents on different technical inventions. Majority from them touch the machines of refrigerations, which a great theorist created in a coauthor with the largest nuclear physicist by Leo Stsylardom, one of creators of the first nuclear reactor.

Compressor constructed by them in 1925 of absorption refrigerator, unlike most widespread presently compressions, did not have locomotive parts (compressor). As khladagenta not freon, but butane or ammonia was used in him. In presented in 1927 on the Leipzig fair the improved refrigerator for cooling a methyl alcohol was used already. In spite of the fact that this refrigerator «howled on the certificate of eyewitnesses, as a jackal», one of electrical engineerings business concerns purchased a patent on him for 12 thousands easily soiled. But in a mass production howling refrigerator so never and was not neglected».

Список использованных источников

1. Энциклопедия «Кругосвет» [Универсальная научно-популярная онлайн-энциклопедия].

URL: http://www.krugosvet.ru/enc/gumanitarnye_nauki/lingvistika/MASHINNI_PEREVOD.html

2. Ляпунов А. А., Кулагина О. С. Использование вычислительных машин для перевода с одного языка на другой. Природа, 1955, № 8.
3. Кузнецов П. С., Ляпунов А. А., Реформатский А. А. Основные проблемы машинного перевода. Вопросы языкознания, 1956, № 5.
4. Панов Д. Ю., Ляпунов А. А., Мухин И. С. Автоматизация перевода с одного языка на другой. В сб.: Сессия по научным проблемам автоматизации производства. М., Изд. АН СССР, 1956.
5. Кулагина О. С. О роли А. А. Ляпунова в развитии работ по машинному переводу в СССР. Проблемы кибернетики, 1977, вып. 32 (в переработанном и дополненном варианте - в книге "Очерки истории информатики в России". Новосибирск, ОИГМ СО РАН, 1998)
6. Кулагина О. С. Исследования по машинному переводу. М., Наука, 1979.
7. Арнольд И. В. «Основы научных исследований в лингвистике» — Высшая школа, 1991 — 140 с.
8. Белоногов Г. Г. «Компьютерная лингвистика и перспективные информационные технологии» —
9. М.: Русский мир, 2004 — 300 с.
10. Марчук Ю. Н. «Проблемы машинного перевода» — М.: Наука, 1983 — 232 с.

УДК 004.896

ПРОЕКТ СТРУКТУРЫ И СХЕМЫ СВЯЗИ ДАННЫХ SMART-UNIVERSITY

Барлыбаев Алибек Бактыбаевич, Акимбекова Эльмира Муханмадиевна

frank-ab@mail.ru

Научные сотрудники НИИ Искусственный интеллект ЕНУ им. Л.Н.Гумилева, Астана,
Казахстан

Научный руководитель – А.А.Шарипбай

В качестве архитектуры смарт-университета была выбрана трехуровневая модель.

Трёхуровневая архитектура – архитектурная модель программного комплекса, предполагающая наличие в нём трёх компонентов: клиентского приложения (обычно называемого «тонким клиентом» или терминалом), сервера приложений, к которому подключено клиентское приложение, и сервера базы данных, с которым работает сервер приложений. На рисунке 1 описана трехуровневая модель смарт-университета.

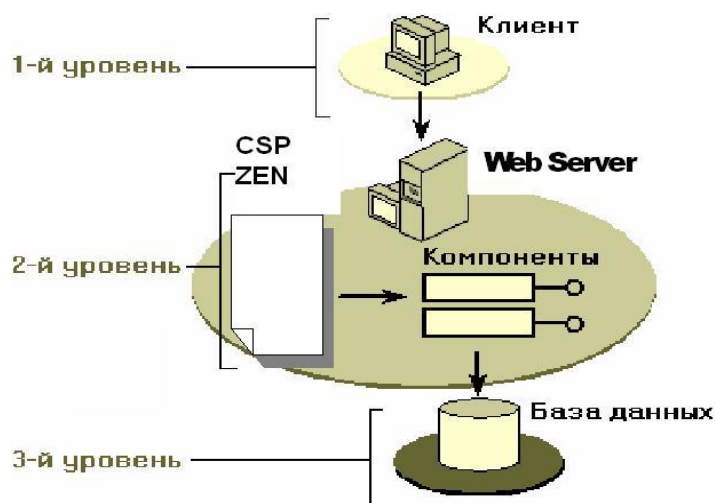


Рисунок 1 – Трёхуровневая модель смарт-университета.