

МИНИСТЕРСТВО ОБРАЗОВАНИЯ И НАУКИ РОССИЙСКОЙ ФЕДЕРАЦИИ
ПЕТРОЗАВОДСКИЙ ГОСУДАРСТВЕННЫЙ УНИВЕРСИТЕТ
ИЖЕВСКИЙ ГОСУДАРСТВЕННЫЙ ТЕХНИЧЕСКИЙ УНИВЕРСИТЕТ
им. М. Т. КАЛАШНИКОВА

Информационные технологии и письменное наследие

El'Manuscript-2012

Материалы IV международной научной конференции
Петрозаводск, 3–8 сентября 2012 года

Петрозаводск, Ижевск
2012

УДК 004.9
ББК 81.11+81.2-0
И741

Изданы при поддержке гранта РФФИ (проект № 12-06-06061-г),
гранта РГНФ (проект № 12-04-14154-г) и в рамках реализации комплекса
мероприятий Программы стратегического развития ПетрГУ на 2012-2016 г.

Ответственные редакторы:

В. А. Баранов, д-р филол. наук, проф.

А. Г. Варфоломеев, канд. физ.-мат. наук, доц.

Информационные технологии и письменное наследие [Текст] :
И741 материалы IV междунар. науч. конф. (Петрозаводск, 3–8 сентября
2012 г.) / отв. ред. В. А. Баранов, А. Г. Варфоломеев. — Петрозаводск ;
Ижевск, 2012. — 328 с.

ISBN 978-5-8021-1402-5

Сборник содержит материалы конференции, посвященной современ-
ным электронным средствам хранения, описания, обработки, исследования
и публикации памятников письменности и исторических источников.

УДК 004.9
ББК 81.11+81.2-0

ISBN 978-5-8021-1402-5

© Петрозаводский государственный
университет, 2012
© Ижевский государственный технический
университет им. М. Т. Калашникова, 2012

КОРПУС ЦЕРКОВНОСЛАВЯНСКИХ ТЕКСТОВ В СОСТАВЕ НКРЯ, ПЕРВАЯ ВЕРСИЯ: ПРОБЛЕМЫ И РЕШЕНИЯ

А. Е. Поляков¹, Е. Р. Добрушина², Т. Ю. Иванова-Алленова²

ИНПБ им. К.Д. Ушинского, Москва;

2Православный Свято-Тихоновский Гуманитарный Университет, Москва

The paper discusses the issues of spelling, grammatical tagging and metatext markup for a pilot version of the Church Slavonic Corpus (started in 2011) and the prospects of its development

1. Корпус церковнославянских текстов¹. За 2011 год группой исследователей² создан и размещен в Интернете (пока идет отладка — на адресе <http://ruscorpora.ru/beta/search-orthlib.html>) пилотный вариант представительного по объему корпуса специальным образом подготовленных церковнославянских текстов. Тексты снабжены специальной разметкой (метатекстовой и грамматической), их общий объем — 1250 текстов, около 4,6 миллиона словоупотреблений, включающих около 150 тыс. различных словоформ. Данный корпус является новой частью проекта "Национальный корпус русского языка" (<http://ruscorpora.ru/>).

2. Тексты корпуса и метатекстовая разметка. Основу корпуса церковнославянских текстов в составе НКРЯ составляют современные богослужебные тексты (XIX–XX вв.) (60%). Кроме того, в корпусе представлены другие периоды (XVII–XVIII вв.) и жанры церковнославянских текстов: писание, святоотеческие, правовые, научные, и др.

Жанрово-тематическая классификация текстов, на основе которой пользователь может создать пользовательского подкорпус, адресована, в первую очередь, массовому пользователю, а не узкому специалисту, и строится исходя из принципов практического удобства и понятности. Метатекстовая разметка, основная концепция которой разработана А.Г. Кравецким, в настоящее время проводится только на уровне целого текста и состоит из следующих параметров (перечисляются корпусные «ярлыки»):

- «Писание» — это Библия, Служебное евангелие и подборки паримий в богослужебных книгах, если они выделены в отдельную рубрику;
- «святоотеческий»;
- «типикон»;
- «служба» — это все богослужебные чины и службы, а также подборки богослужебных текстов (богородичны, кондаки и т. д.) в составе разных сборников;
- «акафист»;
- «право» — это тексты по церковному праву, например, «Книга Правил святых апостолов, святых соборов, вселенских и поместных, и святых отцов»;
- «научный» — например «Ифика Иерополитика».

В будущем предполагается создать метаразметку на уровне отдельных фрагментов текста (например, кондаки, тропари), которые входят в состав больших текстов.

¹ Авторы выражают благодарность за помощь в создании церковнославянского корпуса Программе фундаментальных исследований Президиума РАН «Корпусная лингвистика». Также авторы благодарят за существенную помощь компанию Яндекс, в компьютерном пространстве которой существует весь Национальный корпус русского языка.

² Помимо авторов данных тезисов в создании церковнославянского корпуса на разных этапах его разработки в качестве создателей концепции, разработчиков и исполнителей технических задач, а также консультантов и вдохновителей работы принимали участие В.А.Плунгян, А.Г.Кравецкий, А.И.Зобнин, А.В.Жирова, А.А.Плетнева, Л.И.Маршева, свящ. Ф.Б.Людоговский, Р.Н.Кривко, И.В.Сегалович, свящ. К.О.Польсков.

Также доступен отбор текстов по периоду создания, точнее, по следующим параметрам:

- «архаичный тип» — например, «Добротолюбие»;
- «гибридный тип» — например, «Алфавит Духовный»;
- «стандартный тип» — это все тексты основных богослужебных книг, за исключением текстов XX века;
- «XX век» — это, в первую очередь, акафисты, например «Всем святым, в земле Российской просиявшим».

К сожалению, нет поиска по конкретной дате создания или издания текста, так как определение таких дат для многих текстов требует кропотливой работы, а иногда и вовсе не представляется возможным, более того, для тех текстов, где даты ясны, невозможно определить, в какой момент — в момент создания или в момент издания — производилась последняя языковая правка.

3. Кодировка текста — проблемы орфографии. В процессе подготовки были выработаны принципы и правила представления церковнославянских текстов в корпусе, а также решены практические проблемы, касающиеся разметки и кодировки текста. Каноническая церковнославянская орфография, ориентированная на типографское представление текста, включает много символов, которые фактически не несут смысло-различительной функции. Для представления корпуса в интернете была выработана более простая орфографическая система, которая сохраняет существенные языковые различия, но не пытается имитировать точный типографский вид текста. В этой орфографии отсутствуют некоторые символы с малой или нулевой различительной способностью — придыхания, различие *ou/y*, *ia/я*, — однако сохраняются особенности, связанные с различением лексического значения — *o/omega*, *u/i/v*, *ф/фита*, *з/s*. Кроме того, были выработаны правила перевода текста из этой орфографической системы в современный вид для удобства поиска.

Словоформы в корпусе представлены так же, как в исходном тексте, а леммы даются в унифицированном написании, в котором не используются избыточные буквы, а титла раскрыты (например, слово *млѣть* рассматривается как форма к лемме *милость*).

4. Грамматическое описание. Грамматическое описание церковнославянского словоизменения необходимо для работы лексико-грамматического поиска, а в перспективе и для создания автоматического морфологического анализатора церковнославянского языка.

Существующие грамматики и словари описывают некоторую идеальную картину, которая отражает нормализаторские устремления авторов, но часто не соответствует фактическому состоянию языка. Реально в церковнославянских текстах наблюдается множество орфографических и словоизменительных вариантов, обусловленных тем, что тексты создавались в разное время и в разной языковой среде. Поэтому, если исходить из изучения грамматик и словарей, часто нельзя понять, какие формы имеет некоторое слово. В таких случаях единственным достоверным источником является корпус.

Например, считается, что для отличия мн. и дв. числа от омонимичных форм ед. числа используется обложенное ударение (^) и замена $o \rightarrow w$, $e \rightarrow \epsilon$. Спрашивается, как будет им. мн. от *высота* — *высоты^*, *высшты'* или *высшты^*? Оказывается, в корпусе встречаются все три формы, причем первая 9 раз, вторая — 17, третья — 1. Из анализа других аналогичных случаев мы можем сделать вывод, что при конкуренции правил замена $o \rightarrow w$ предпочтительнее, чем обложенное ударение.

К настоящему моменту была проделана значительная работа, результатом которой явилась пилотная версия грамматического словаря. Было проанализировано около 150 тыс. словоформ, большинству из них была приписана основная грамматическая информация: лемма, часть речи, грамматические признаки. Вначале была сделана ручная лемматизация для ряда словоформ, а затем на основе построенной модели словоизменения были проанализированы остальные словоформы. Часть словоформ была квалифицирована как ошибки (опечатки). Для некоторых редких форм лемму определить не удалось или она была определена гипотетически. Поиск по переменным характеристикам, таким как падеж или род (для прилагательного), работает, но без снятия грамматической омонимии.

Грамматические таблицы, на которые опирается автоматический анализатор, строятся на основе анализа корпуса текстов и существующих грамматик. Поскольку грамматики дают неполную и противоречивую картину словоизменения, то только корпус является окончательным критерием истины.

В отличие от традиционных грамматик, парадигмы не были заданы априорно, а выведены эмпирически на основе анализа множества словоформ, имеющих однотипное соотношение между грамматическими формами. Таким образом, номенклатура парадигм получается значительно более детальной, чем традиционная, и подчас существенно от таковой отличается. Например, традиционное первое склонение (*рабъ*) на самом деле распадается на 14 подтипов в зависимости от конечного согласного (парный твердый — мягкий, велярный, шипящий, йот), наличия беглого гласного и других особенностей (*іерей* — *іереа*, *іереомъ*, *агарянинъ* — *агаряне*).

5. Перспективы церковнославянского корпуса. Чтобы корпус мог в полной мере удовлетворять потребностям современной науки и образования в области исследования и преподавания церковнославянского языка, требуется значительная доработка и развитие. Хотелось бы в будущем сделать следующее:

- Снабдить корпус справочными материалами об особенностях корпуса и составе текстов, облегчающими пользователям работу с корпусом.
- Снабдить корпус кратким словарем, поясняющим термины, использованные в разметке текстов.
- Выверить грамматические характеристики словоформ, приписанные компьютерными методами, и устранить ошибки.
- Приписать леммы и грамматические характеристики словоформам, оставшимся неразобранными в результате программного анализа.

- Произвести путем анализа необработанных слов и их значений в текстах поиск ошибок, возникших при наборе текстов, и выправить эти ошибки.
- Проанализировать значение лемм, определенных разметчиками как имена нарицательные, и приписать тем из них, перевод которых на русский язык не ясен без привлечения специальных знаний, краткое толкование, доступное пользователю вместе с грамматической характеристикой при нажатии на выбранное слово на странице результатов поиска.
- Проанализировать значение лемм, определенных разметчиками как имена собственные, подтвердить или отвергнуть это решение и определить для каждого имени собственного лексико-семантический класс: имя, топоним и др. Такая работа позволит расширить представления о лексическом составе церковнославянского языка и употребляющихся в нем имен собственных, даст возможность ликвидировать ошибки массовой разметки, а также даст возможность пользователям легче понимать содержание полученного текста.
- Снабдить корпус списком основных лемм и создать для пользователя возможность переходить к поиску интересующей его леммы непосредственно из алфавитного списка, что во многом решит проблемы, связанные с вариативностью орфографии и наличием в запросах букв, требующих специальных шрифтов или использования виртуальной клавиатуры.