

Разработка системы для сбора данных из открытых источников

Лебедев Станислав Филиппович

руководитель от предприятия: В. Д. Танков

руководитель от факультета: А. А. Шпильман

Санкт-Петербургская школа
физико-математических и компьютерных наук
НИУ ВШЭ – Санкт-Петербург

28 июня 2021 г.

Команда “Grazie” компании JetBrains решает задачи обработки естественного языка:

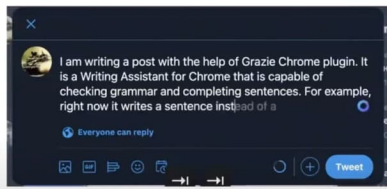
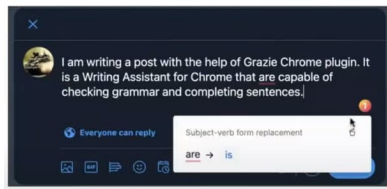
- Автодополнение
- Суммаризация
- Перефразирование
- Исправление грамматических ошибок

Для этого требуется много данных, открытые источники – отличное решение для их сбора:

- В интернете много разной информации. Можно набирать датасеты как узкоспециализированные, так и “общего” плана
- Открытые источники помогают в коммерциализации продукта

Задача исправление грамматических ошибок (GEC)

Задача исправления грамматических ошибок –
детекция и исправления грамматической ошибки в
тексте на естественном языке



Подход схож с описанным в статье: GECToR – Grammatical Error Correction: Tag, Not Rewrite Kostiantyn Omelianchuk, Vitaliy Atrasevych, Artem Chernodub, Oleksandr Skurzhashnyi Grammarly 15th Workshop on Innovative Use of NLP for Building Educational Applications (co-located with ACL 2020) [<https://arxiv.org/pdf/2005.12592.pdf>]

Особенности существующих решений

- Некоммерческие. Невозможно создавать коммерческий продукт.
- Сложно и дорого расширять, так как требуется работа живых людей.
- Смещение в область художественной литературы.
- Недостаточно объемны

Схема сбора данных

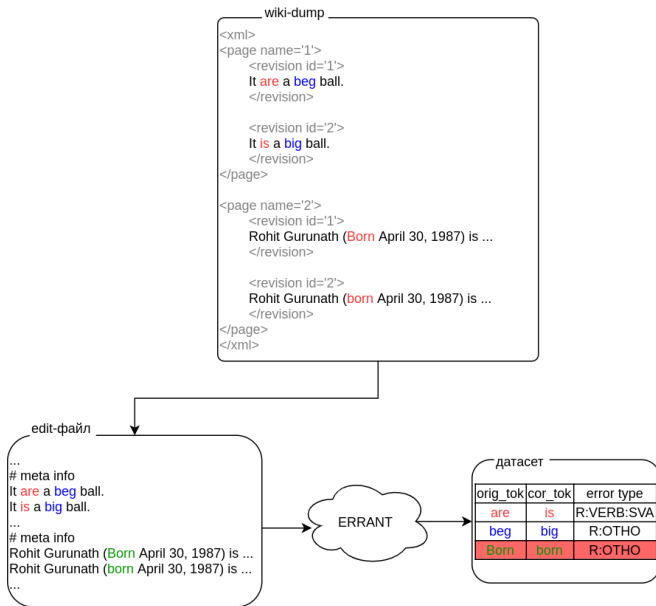


Ревизии – сохраненное состояние страницы

Дампы – архивированные наборы ревизий

Edit-файлы – пары предложений с внесенными изменениями

ERRANT – фреймворк для классификации изменений



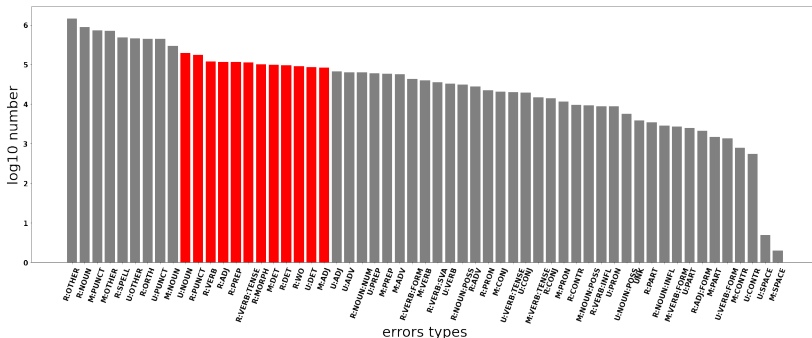
Цель: реализовать инструмент извлечения датасета из Википедии для задачи исправления грамматических ошибок.

Задачи:

- Изучить дампы и выдвинуть гипотезы о методах фильтрации
- Разработать программу для извлечения датасетов в соответствии с гипотезами
- Собрать датасеты и проверить их качество

Фильтрация по индексу авторов

- Возьмем авторов с больше чем n ревидиями.
- Индекс – доля ревидий, попавших в класс “хороших”

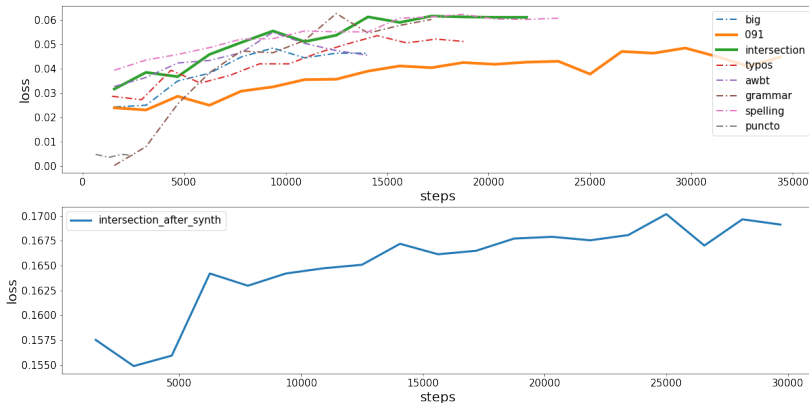


Фильтрация производилась регулярными выражениями.

- Отдельные выражения по упоминанию грамматики, написания и опечаток
- Комбинированный поиск
- Проект AutoWikiBrowser – проект на Википедии, собирающий и валидирующий регулярные выражения для поиска грамматических ошибок. В своих ревизиях оставляют специальный тэг

Валидация 1

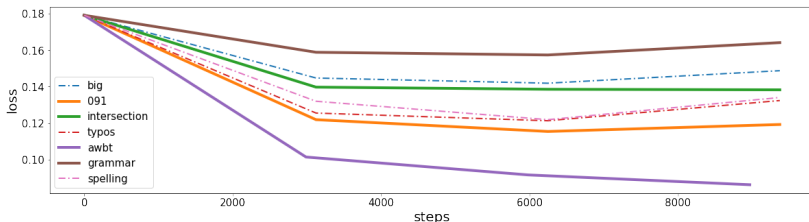
Датасет на 100 000 примеров



- Фильтрация по комментариям работает лучше, чем фильтрация по индексу
- Фильтрация по комментариям улучшает синтетику

Валидация 2

Датасет на 200 000 примеров



Из-за улучшения сентетики все датасеты показали себя хуже, однако можно сделать и положительные выводы:

- **Фильтрация по комментариям** работает лучше, чем **фильтрация по индексу**
- **Фильтрация по комментариям** иногда почти не портит сентетику

Из странное: **фильтрация по комментариям проекта AutoWikiBrowser** показывает себя очень плохо

- Разработана система создания датасетов от списка ссылок на дампы до готовых к запуску датасетов
- Изучены разные подходы фильтрации датасетов
- Обнаружено что фильтрация по комментариям работает эффективнее

- NUCLE
- Lang-8
- The First Certificate in English
- WI+LOCNESS

Доступно к скачиванию на
<https://www.cl.cam.ac.uk/research/nl/bea2019st/>