

ФОРМАЛЬНИЙ АНАЛІЗ МОРФЕМНИХ СКЛАДОВИХ ТЕЗАУРУСА МОВНИХ ОБРАЗІВ

Кравчук І.А.

Науковий керівник – проф., д.т.н. Кветний Р.Н.

Лінгвістична обробка природномовних (ПМ) текстів є однією з центральних проблем інтелектуалізації інформаційних технологій. Традиційно аналіз ПМ текстів представлявся як послідовність процесів – морфологічний аналіз, синтаксичний аналіз, семантичний аналіз.

Морфологічний аналіз вхідних слів є першим етапом аналізу ПМ тексту. Результати морфологічного аналізу є базовою основою для всіх наступних етапів автоматичного аналізу та синтезу тексту. Морфеми поділяються на два основних типи – кореневі та афіксальні. Корінь є основною значущою частиною слова. Афікс є допоміжною частиною слова, що приєднується до кореня і служить для словотвору та вираження граматичних значень. Найбільш поширені два типи афіксів – префікси, що знаходяться перед коренем, та постфікси, що знаходяться після кореня. Постфікси розділяються на суфікси (мають словотвірне значення) і флексії або закінчення (вказують на зв'язок з іншими членами речення). В українській мові, яка належить до флексійних мов, використовуються як префікси, так і постфікси (суфікси та закінчення).

Аналізуючи морфологічну структуру слова флексійних мов, визначимо базові правила з P для слів w та u :

- а) x є префіксом w ($x \triangleright w$), якщо $w = xu$,
- б) z є суфіксом w ($z \triangleleft w$), якщо $w = uz$,
- в) v є закінченням w ($v \triangleleft w$), якщо $w = uv$ або $w = uzv$.

В загальному випадку слово з одним префіксом, суфіксом та закінченням має вигляд:

$$w = xuzv.$$

Формалізуємо правила визначення множин коренів, префіксів та суфіксів:

$$Root(L) = \{y \mid (\exists u \in L)\},$$

$$Pref(L) = \{x \mid (\exists u \in L)x \triangleright w\},$$

$$Suf(L) = \{z \mid (\exists u \in L)z \triangleleft w\},$$

$$End(L) = \{v \mid (\exists u \in L)v \triangleleft w\}.$$

За допомогою цих правил забезпечується можливість перевірки належності довільного слова w до множини $L(G)$ всіх правильних слів у граматиці, що розглядається.

Отже, задача зводиться до обґрунтування конструктивних алгоритмів попереднього відокремлення морфем з складових тезауруса мовних образів.