

Найбільшу кількість варіантів у нашому дослідженні мало прислів'я, яке не одержало жодного голосу на підтримку традиційного варіанту: *Haléř k haléři, a ze dvanácti bude groš*. Одним із пояснень цього факту може бути наявність у його складі назви архаїчної реалії *groš* (від лат. *grossus* – "товста монета") – монети, що в Середні віки була срібною валютою, а в XVII-XIX ст. – зовсім дрібною платіжною одиницею, пор. її оцінні ознаки: *bídny, žebrácký, hladový, těžce vydělaný, poslední groš, pár grošů* [4, s.552]. Наші інформанти в основному правильно оцінили смисл вислову, що полягав в ідеї економного ставлення до грошей, і запропонували такі продовження прислів'я: *Haléř k haléři, a koruna ke koruně // korunka ke korunce = 16; Haléř k haléři, a koruna je z toho = 1; Haléř k haléři, a už mám dva = 1; Haléř k haléři, a koruna je celá = 1; Haléř k haléři, a koruna je doma = 1; Haléř k haléři, a je koruna = 1; Haléř k haléři, a je celá koruna = 1; Haléř k haléři, a koruna je hned = 1; Haléř k haléři, a truhlička je plná = 1; Haléř k haléři, a budeš boháč = 1; Haléř k haléři, a bude korunka // Babka k babce budou kapce = 1; Haléř k haléři, a máš korunu = 1; Haléř k haléři, a je z toho milion = 2; Haléř k haléři, a je na kravu = 1; Haléř k haléři, a prase je celé = 1.*

Варіант *Co není v hlavě, musí být v nohách* (укр. *За дурною головою і ногам нема спокою*) цікавий тим, що незалежно від віку 5 інформантів утворили контамінований вислів: *Co není v hlavě, musí být v nohách // Komu není z hůry dáno, v apatyce nekoupí → Co není v hlavě, v apatyce nekoupíš*.

Результати нашого анкетування засвідчили, поперше, високий ступінь варіантності паремійних одиниць. У всіх інформантів зі 100 прислів'їв 84 одиниці одержали 885 різних закінчень. По-друге, молоде покоління, не знаючи традиційного варіанта прислів'я, дописує таке його закінчення, яке не має ніякого відношення до паремії як такої, а радше є асоціативною реакцією на буквально прочитання її першої частини, пор.: *Co se doma uvaří, tak já sním* (13 let); *Jak si usteleš, tak si vyspíš = 5* відповідей; *Boží mlýny melou mouku* (13 let); *Dobrá hospodyňka se dobře vdá* (17 let); *Jedna ovce prašívá mě včera potkala na nadraží* (16 let); *Kdo chce kam, ať jde tam* (25 let); *Kdo chce s vlky býti, musí umět výt* (14 let); *Když se kácí les, ubírá to kyslíkové zdroje* (17 let); *Komu není z hůry dáno, tak ten má smůlu* (17 let);

Neštěstí nechodí, neštěstí se děje (17 let); *Kdo chce psa bít, ať si opatří hůl* (20 let); *Kdo šetří, má peníze* (25 let); *Kdo rychle dává, brzo nemá nic* (16 let); *Kdo rychle dává, bude brzo chudej* (16 let) та ін.

Деякі варіанти відповідей позначено довільною мовною творчістю респондентів, пор.: *Líná huba mele vodu* (17 let); *Mezi slepými vidomého nehledej* (17 let); *Neštěstí nechodí pěšky* (14 let); тим більше нас здивувала відповідь *Pečení holubi lítají do huby* 4 інформантів середнього віку (35, 34, 50 років і один без зазначення віку).

Жодного варіанта не мали 15 прислів'їв: це саме ті одиниці, що увійшли в паремійний мінімум чеської мови, а саме: *Bez práce nejsou koláče; Dvakrát měř, jednou řež; Jablko nepadá daleko od stromu; Jak se do lesa volá, tak se z lesa ozývá; Kdo maže, ten jede. Kdo pozdě chodí, sám sobě škodí; Komu není rady, tomu není pomoci; Lež má krátké nohy; Neříkej hop, dokud nepřeskočíš; Pes, který štěká, nekouše; Pro dobrotu na žebroutu; Šaty dělají člověka; Tichá voda břehy mele; V nouzi poznáš přítele; Žádný učený z nebe nespadl*.

Аналіз цих одиниць щодо їх походження свідчить, що всі вони відзначаються ґрунтовним історичним стажем, а отже, викликає сумнів твердження Д. Біттнерової і Ф. Шиндлера, "що збірки Челаковського і Спілки є вкрай застарілими" [2, s. 297].

Загальний висновок, що випливає з нашого дослідження й дослідження Д. Біттнерової й Ф. Шиндлера, полягає в тому, що незважаючи на велику різницю в кількісному складі аналізованих паремійних одиниць (100 проти 11 000), результат виявився майже однаковим: половина прислів'їв і приказок, рівень знання яких інформантами становив понад 95%, повністю збігається.

Перспективи подальшого дослідження цієї проблеми ми вбачаємо в залученні до дослідження даних Чеського національного корпусу, що послужило б підґрунтям для порівняльного аналізу паремійних мінімумів в усному і писемному дискурсах.

1. Пермяков Г.Л. Основы структурной паремииологии. – М., 1988; 2. Bittnerová D., Schindler F. Česká přísloví. Soudobý stav konce 20. století. – Praha, 1997; 3. Čelakovský F.L. Mudrosloví národu slovanského ve příslovích. – Praha, 1949; 4. Slovník spisovného jazyka českého / Zpracoval lexikografický kolektiv Ústavu pro jazyk český ČSAV pod vedením Boh. Havránka. – Praha, 1971. – D. 1.

Н. Дарчук, канд. філол. наук

ДОСЛІДНИЦЬКИЙ КОРПУС УКРАЇНСЬКОЇ МОВИ: ОСНОВНІ ЗАСАДИ І ПЕРСПЕКТИВИ

Запропоновано методику створення корпусу українських текстів і технологію його конструювання в режимі on-line, що допоможе ефективно й оперативно здійснювати масштабні комплексні філологічні дослідження на рівні сучасної наукометрії. Створений корпус може використовуватися як готовий продукт і як модель для конструювання аналогічних корпусів колегами з різних наукових і навчальних закладів України. Результати викладені в мережі Інтернет для загального користування на лінгвістичному порталі: www.mova.info/corpus.aspx.

The article describes the method to produce the corpus of the Ukrainian language and the technology for its on-line mode construction which enables to conduct efficient large-scale and comprehensive philological studies and research works on the level of modern sciencemetrics. The corpus can be used as final product as well as a model for construction of similar corpora by specialists from academic and education institutions of Ukraine. The project results can be found at the linguistic web-site www.mova.info/corpus.aspx for common use.

Стан лінгвістики останніх десятиліть Ю. Апресян [1] визначає метафорою "золотий вік лексикографії" через безпрецедентне зростання числа й різноманіття словників, удосконалення методики лексикографування, принципів системного опису лексики й увагу сучасної теоретичної лінгвістики до опису окремих лексем. Такий прорив у мікросвіт значення одного слова спонукає філологів до всебічного лінгвістичного "портретування", а значить інтегрального опису мови, що аж ніяк не можливо без підтримки комп'ютерної лінгвістики в цілому і корпусної зокрема.

Справді, фонетичні, просодичні, морфологічні, синтаксичні, семантичні, прагматичні, стилістичні, комунікативні, сполучуванісні властивості лексем є вивідними

безпосередньо з текстів, які були і є невичерпним джерелом для словників і граматики.

Незаперечним є й той факт, що в останні десятиліття фахівці галузей, що мають справу з комп'ютерним аналізом текстів, гостро відчують потребу в якомога більшому числі функціональних характеристик мовних одиниць для різних типів текстів. Для проведення теоретичних і прикладних (машинний переклад, аналіз і синтез мовлення, анотування і реферування тексту тощо) та дидактичних досліджень спеціалісти в галузях когнітивної лінгвістики, стилістики, семасіології, функціональної граматики, словотвору не мають узагальнених даних про закономірності функціонування мовних одиниць у мові

ленні. Усю необхідну лінгвістичну інформацію для подальшого опрацювання її у філологічних студіях можна одержати з корпусів текстів, розбудова яких є ознакою нашого часу.

В останні десятиліття минулого сторіччя зусилля багатьох країн були спрямовані на створення національних універсальних корпусів текстів із необхідними й достатніми кількісними та якісними параметрами для укладання на їхній основі словників і граматик національних мов. На сьогодні текстові корпусні ресурси вже існують для багатьох мов, і не лише європейських, створення корпусу вважається за обов'язок по відношенню до національної мови. Серед найвідоміших – Британський національний корпус (100 млн. слововживань); Національний корпус російської мови (10 млн. слововживань); Корпус австралійської періодики (300 млн. слововживань); Банк англійської мови (320 млн. слововживань); Корпус німецької мови (778 млн. слововживань) тощо, а в маленькій Словаччині існує навіть Інститут корпусного дослідження словацької мови.

Проте українська мова все ще залишається однією з небагатьох, котра не має репрезентативного корпусу, мета якого – відобразити як сучасний, так і історичний стан розвитку та функціонування української мови. В українському мовознавстві не розвивається напрямок корпусної лінгвістики, мета якої – фіксувати всі аспекти функціонування мови, зберігаючи як інтра-, так і екстралінгвістичну специфіку й адекватність лінгвістичних даних, а завдання – розробити принципи і методи побудови та експлуатації текстових корпусних ресурсів із використанням сучасних комп'ютерних технологій.

Існує кілька шляхів розвитку корпусної лінгвістики в Україні, але спільним і неодмінним для них є накопичення текстів різних стилів і часових зрізів із відповідним зовнішнім (джерело, автор, рік та місце видання тощо) і внутрішнім, пов'язаним зі структурацією тексту, маркуванням (номер розділу, абзацу, речення, слова, початок і кінець тексту, абзацу, речення тощо).

Перший шлях полягає в програмному забезпеченні пошуку окремої словоформи в певному джерелі або в усьому корпусі текстів і формуванні для неї дослідницької ілюстративної текстової бази. Така система може функціонувати без використання лінгвістичних алгоритмів і словників або з використанням словників як допоміжного матеріалу для формування запитів, хоча сам корпус лінгвістично не розмічається.

Другий шлях пов'язаний з розробленням програмного забезпечення для маркування корпусу текстів самим дослідником. За допомогою спеціальних програм та інтерфейсу користувач-дослідник може "портретувати" слово: визначати його морфологічну, синтаксичну, лексичну характеристику згідно з контекстом уживання. Цей шлях забезпечує можливість, з одного боку, одержання однозначної, точної в лінгвістичному смислі інформації, з іншого – створення лексикографічної картки за результатами розмітки.

Третій шлях передбачає наявність комп'ютерних інструментів – пакету програм для різнобічного лінгвістичного опрацювання текстового матеріалу. У цьому разі корпус текстів лінгвістично розмічається в автоматичному режимі з подальшою автоматизованою обробкою для зняття омонімії. Завдяки цьому користувач одержує інформацію, передбачену програмним забезпеченням та базами даних (словниками з морфологічною, синтаксичною, лексичною, енциклопедичною інформацією), створює за текстами свої власні алфавітні й частотні словники, робить автоматично транслітерацію, фонематичний або фонетичний запис, сегментує текст на морфи, морфологічно індексує текст / слово, здійснює синтаксичний аналіз речення тощо.

Кожен із цих шляхів має свої позитивні риси й недоліки. До **позитивних** рис першого шляху можна віднести темпи створення текстових корпусів; другого – точність характеристики одиниць аналізу та при узгодженій роботі спільними зусиллями фахівців – створення сучасної автоматизованої лексикографічної картотеки; третього – ефективність і швидкість в одержанні різноманітної інформації. До **недоліків** – ресурсну обмеженість першого, повільність роботи другого та недостатню точність третього, тому що 100% бездоганного опрацювання тексту автоматично без постредагування досягти практично неможливо.

Займаючись проблемами корпусної лінгвістики, співробітники лабораторії комп'ютерної лінгвістики Інституту філології обрали третій шлях, розробивши відповідну концепцію, суть якої полягає в анотуванні текстових корпусів за якісними і кількісними ознаками різних мовних одиниць, у розробленні методичних засад створення комп'ютерних інструментів для лінгвістів, а також пакетів програм для укладання електронних картотек, алфавітно-частотних словників, тезаурусів, словників морфемно-словотворчих гнізд тощо. Така засаднича база розвиває проект корпусного опрацювання вшир і вглиб.

Вироблена класифікація корпусів текстів має таку типологію: **національні / дослідницькі, ілюстративні, паралельних текстів, динамічні і статичні**. На даний час наш корпус текстів є **дослідницьким**, оскільки призначений для вивчення різних аспектів функціонування мовної системи і зорієнтований на широкий клас лінгвістичних завдань; **статичним**, оскільки відображає певний часовий стан мовної системи. Наприклад, авторські корпуси поетичної мови – це колекції текстів окремих поетів кінця ХХ ст., публіцистичні тексти репрезентують мовлення газет за короткий проміжок часу (2008–2010 рр.). Водночас ці тексти дають можливість дослідити мовні явища в динаміці, тому він є **динамічним**, а також **інтерактивним**, оскільки користувач, здійснюючи дослідження, може виділити з генерального корпусу робочий корпус, який включає лише частину генерального корпусу, наприклад тексти якоїсь тематичної зони (Політика, Економіка, Культура, Спорт) або якоїсь газети ("Дзеркало тижня", "Україна молода", "Українська правда", "Високий замок"). Накопичено також корпус **паралельних українсько-російських текстів** (публіцистичний стиль – по 500 тис. слововживань для обох мов; художній стиль – по 300 тис. слововживань).

Не зупиняючись спеціально на загальних проблемах і вимогах до корпусу текстів, зауважимо, що добір текстового матеріалу, його репрезентативність, структуризація, організація за стилями, зонами, хронологією є для нас обов'язковими умовами організації корпусної системи української мови.

Головна увага приділяється анотуванню корпусного матеріалу, тобто процесу введення формалізованої лінгвістичної інформації в електронний текст. Аплікативне призначення корпусних даних – фонологічні, морфологічні, синтаксичні, лексикологічні, лексикографічні тощо дослідження – детермінує тип лінгвістичної анотації корпусу. За умови використання корпусного ресурсу в дослідженнях **фонемного / фонетичного аспектів** мови обов'язково має бути фонетична та просодична анотації корпусу. **Морфологічна анотація** передбачає визначення морфологічних параметрів слова: частиномовну приналежність і категоріальні ознаки кожної словоформи тексту. **Морфемна анотація** полягає в сегментуванні кожної словоформи тексту за типами морфів і подальше автоматичне укладання серії морфемних та словотвірних словників із частотними характеристиками морфа/морфемами, за якими

можна вивчати комбінаторно-дистрибутивну будову. **Синтаксична анотація** пов'язана з автоматичним обробленням кожного речення і виділенням словосполучень та приписуванням кожному з них інформації про тип (дієслівний, іменниковий тощо), вид синтаксичного зв'язку та семантичного відношення. Для кожного речення автоматизовано представлятиметься синтаксична структура на рівні членів речення. **Лексико-семантична інформація**, яка буде приписана кожному слову в тексті, складатиметься з трьох частин:

Розряд (напр., загальна, власна назва, займенник-іменник, займенник-прикметник тощо).

Власне лексико-семантичні характеристики (тип за таксономічною класифікацією)

Напр., Іменники предметні

Особа, в тому числі:

Етноніми

Назви членів родини

Надприродні істоти

Тварини

Рослини

Речовини, матеріали

Простір і місце

Будівлі і споруди

Інструменти і пристрої

Інструменти

Механізми і прилади

Транспортні засоби

Зброя

Музичні інструменти

Меблі

Посуд

Одяг, взуття

Їжа, напої

Дериваційні (словотворчі) характеристики (напр., "демінітив", "відад'єктивний прислівник" тощо).

Мета нашого проекту полягає в комплексній та аспекти розробці корпусу текстів таких підкорпусів (кількісні дані наводяться за сучасним станом проекту):

- української поезії (кінець XX ст.) (близько 1 млн.);
- української художньої прози (кінець XX ст.) (300 тис.);
- фольклорних текстів (32 тис. слововживань);
- публіцистичних текстів за 2008–2010 роки (2 млн.);
- наукових текстів з лінгвістики, економіки, екології, правничих, навчальних текстів (150 тис.);

• паралельних українсько-російських текстів (близько 1 млн. слововживань).

Наш корпус української мови побудований таким чином, щоб користувач міг одержати різноаспектну інформацію про мовну одиницю (морфема, слово, словосполучення, речення):

- **зовнішню** – анування тексту (автор, джерело, рік видання тощо);

- **внутрішню** – структурування тексту (номер розділу, абзацу, речення, слова тощо), його **жанрові особливості й тип тексту**.

Наприклад, **художні тексти** маркуються за жанровими позначками: *історико-пригодницький, кримінальний, сатиричний і гумористичний, фантастичний*

*тексти тощо; тип тексту: анекдот, детектив, оповідання, роман, повість, казка тощо. Якщо відомий **хронотип тексту**, особливо для художніх текстів як інформативна характеристика (місце і час описуваних подій, напр., Україна XIX ст. або Україна/революція, Україна/війна 1941–1945 рр., або Україна/пострадянський період тощо), то він вказується.*

Нехудожні тексти супроводжуються вказівкою на сферу функціонування тексту (*побутова, офіційно-ділова, виробничо-технічна, публіцистична, навчально-наукова, церковно-богословська*), що пов'язано з мовленнєвими особливостями тексту. Визначаються такі **типи нехудожніх текстів**: *автобіографія, щоденник, довідник, договір, інструкція, кодекс, коментар, лист оголошення, резюме, твір, хроніка тощо.*

Публіцистичні тексти обов'язково маркуються за **тематикою тексту** (*бізнес, комерція, економіка, фінанси; війна і військові конфлікти; дім і побут; здоров'я і медицина; мистецтво, наука, політика і суспільне життя, право, виробництво, сільське господарство, спорт, природа тощо*). Таке метатекстове розміщення є міжнародним (стандарт EAGLE Дж. Сінклера з додатками для слов'янських текстів С. Шарова), покликаним полегшити в майбутньому зіставлення результатів у міжнародному Корпусі.

- граматичну, лексико-граматичну, лексичну інформацію про мовні одиниці;

- кількісну – про функціонування її як інваріантної форми (лема, морфема) і варіантної (словоформа, морф) зі статистичними характеристиками (наразі тільки абсолютна частота, а в перспективі – середня частота, міра коливання середньої частоти, коефіцієнт стабільності) (див. рис. 1, 2, 3).

Така сукупність характеристик дасть можливість створювати **параметризовану базу даних**, під якою ми розуміємо багатоаспектну і багатофункціональну систему, яка включає:

- 1) корпус текстів – джерело різного роду словників;
- 2) серію алфавітно-частотних словників з усією інформацією про слово: граматичною, лексико-граматичною, стилістичною та статистичною (абсолютна та середня частота, міра коливання середньої частоти, коефіцієнт стабільності);
- 3) алфавітно-частотні словники слів та слововживань спільної лексики.
- 4) словники неолексем;
- 5) словники синтаксичних моделей керування: дієслівних, іменникових, атрибутивних;
- 6) серію морфемних та словотвірних словників із частотними характеристиками морфа/морфеми, за якими можна вивчати комбінаторно-дистрибутивну будову, словотвірне значення кожної афіксальної морфеми в текстах;
- 7) словники синонімів, антонімів, фразеологізмів, тезауруси;
- 8) словник-конкорданс як допоміжний інструмент для формування лексико-семантичної, синтаксичної та стилістичної характеристики слова.

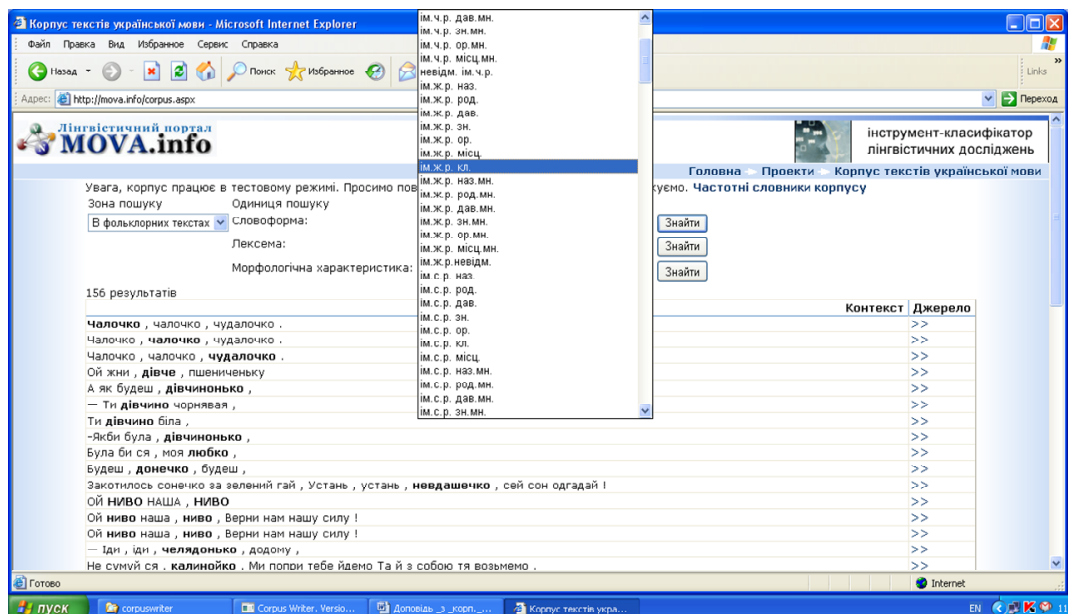


Рис. 1. Інтерфейс інтернет-програми з побудови конкордансу до морфологічної інформації

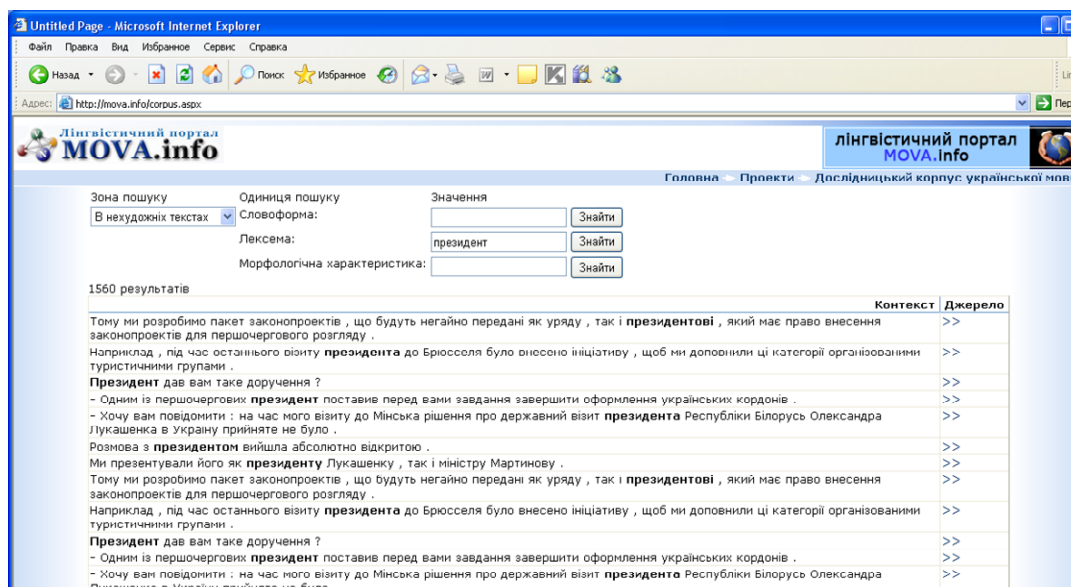


Рис. 2. Інтерфейс інтернет-програми з побудови конкордансу до лексики/словоформи

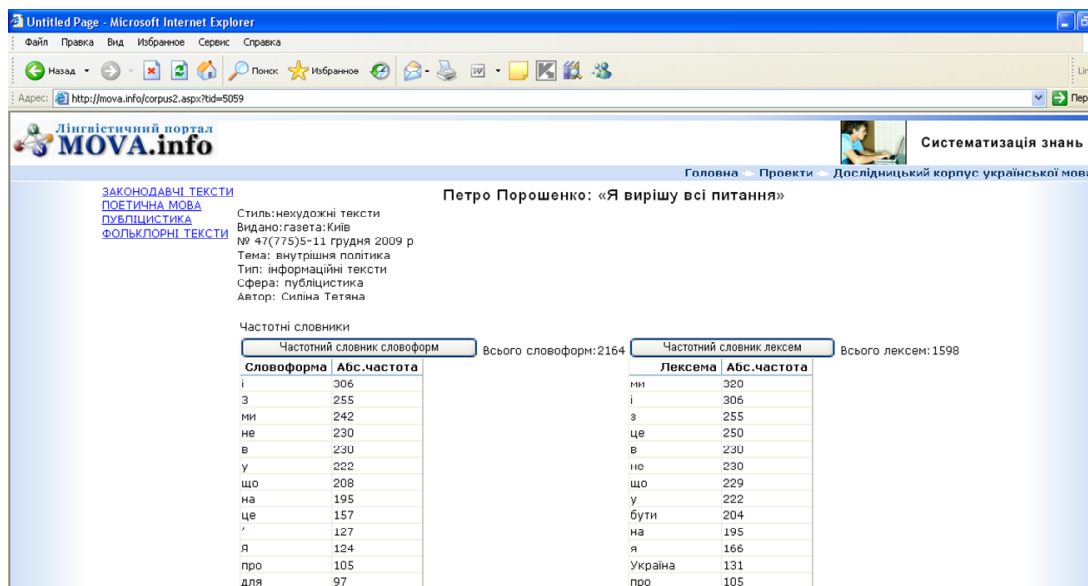


Рис. 3. Інтерфейс інтернет-програми з побудови алфавітно-частотного словника до конкретного тексту

Ми розглядаємо параметризовану базу як реальну можливість для створення універсального (багатоцільового) словника. У базі передбачено параметри:

- граматичні (частина мови і категоріальні значення, напр., рід, число, відмінки, особа тощо);
- структурні (напр., моделі морфної структури слів різних частин мови; моделі керування дієслівні, іменні, атрибутивні тощо);
- лексико-семантичні, які віддзеркалюють системні відношення (синонімія, антонімія, омонімія, паронімія, ідеографія);
- статистичні.

Таке багатоаспектне накопичення мовних фактів та встановлення закономірностей їхнього функціонування в різних типах текстів разом із фреквентарієм дозволить розробити, за словами М. В. Всеволодової, новий тип граматики – граматики для лінгвіста. І перші кроки в цьому напрямку ми вже можемо робити сьогодні, розробляючи такі проблеми:

Системно-структурне дослідження морфології української мови:

для іменників:

- динаміка використання різних відмінків і тенденції розвитку відмінкової системи;
- сучасні тенденції в родовій категоризації іменника;
- словозмінні та класифікаційні елементи в категорії числа (семантика злічуваності / незлічуваності);

для прикметників:

- морфолого-синтаксична характеристика прикметника;
- морфологічні категорії прикметника;
- перехід інших частин мови в прикметники тощо;

для дієслова:

- обсяг реалізації дієслівної парадигми та особливості використання дієслівних форм у сучасних дискурсах;
- категорія виду та видові протиставлення в сучасних текстах;
- категорія часу та явища транспозиції тощо.

Комп'ютерна підтримка проекту забезпечується пакетом алгоритмів і програм автоматичного морфологічного і

синтаксичного аналізу української мови, морфного сегментування тексту, укладання повних і часткових словників за частотою, алфавітом, а також пакетом статистичної обробки лінгвістичних даних. Створення алфавітного словника здійснюється у двох режимах: автоматичному (без зняття омонімії) і автоматизованому (зі зняттям лексичної і лексико-граматичної омонімії, у деяких випадках ця функція також передбачена пакетом програм). Для виконання таких завдань, як синтаксичний аналіз речення, створення словників неолексем, синонімічних груп, тезауруса тощо передбачено програмне забезпечення в автоматизованому режимі, оскільки їх здійснення потребує контроль за значенням. При цьому обов'язковою є функція конкордансної ілюстрації з корпусу текстів. Програмне забезпечення базується на таких словниках української мови: орфографічний, граматичний, морфемний, синонімічний, антонімів, фразеологічний, тлумачний, кожен з яких має комп'ютерне представлення у базі знань, у розробці якого брали участь автори цього проекту. Програмно передбачено автоматичне порівняння авторських текстових словників зі словниками мови, що скорочує роботу лінгвіста в десятки разів.

Отже, запропонована методика створення корпусу українських текстів є узагальненням комплексу теоретичних і прикладних ідей сучасного мовознавства. Технологія конструювання корпусу робить її надзвичайно ефективним та раціональним інструментом (вона зберігає час та людські ресурси) для спеціалістів-філологів різного профілю, тому що передбачає роботу в режимі **on-line**. Електронна база лінгвістичних даних – компонент корпусного дослідження – зі своєю методологією і технологією допомагає ефективно й оперативно здійснювати масштабні комплексні філологічні дослідження на рівні сучасної наукометрії.

Створений корпус може використовуватися як готовий продукт і як модель для конструювання аналогічних корпусів колегами з різних наукових та навчальних інституцій України.

1. Апресян Ю. Д. О толковом словаре управления и сочетаемости русского глагола // Словарь. Грамматика. Текст. – М., 1998.

Н. Дем'яненко, канд. філол. наук

ОСОБЛИВОСТІ СЕМАНТИЧНОГО СПІВВІДНОШЕННЯ РОЗУМУ І ПОЧУТТІВ У ПОЛЬСЬКІЙ ФРАЗЕОЛОГІЇ

Проаналізовано групу фразеологізмів польської мови в руслі проблематики картини світу, досліджено особливості семантичного співвіднесення розуму і почуттів, представлених у польській фразеології.

The article analyzes a number of Polish idioms within language pictures of the world, specific features of semantic correlation of rationality and feelings represented in the Polish phraseology.

Розум і почуття відіграють важливу роль у створенні характеристики людини. Пізнавальні потреби передбачають появу інтелектуальних рис характеру, потреба в спілкуванні обумовила появу емоційних рис характеру. У польській фразеології знайшли своє відображення як одні, так й інші риси характеру. Дослідження фразеологічної системи польської мови дозволили виявити, що найбільша кількість фразеологічних одиниць на позначення розумових здібностей людини містить у своєму складі лексему "голова". Найбільш частотною серед фразеологізмів, що окреслюють почуття людини, є група з головним компонентом "серце". Отже, перед нами стоїть завдання дослідити фразеологічні одиниці з головними компонентами "голова" і "серце" в польській мові. Дослідники наголошують на тому, що назви частин тіла використовуються при створенні метонімічних сполучень слів, доповнюючи й збагачуючи номінацію і, таким чином, є визначним фактором розвитку лексичної системи будь-якої мови. Ролі компонентів у семантичній структурі фра-

зеологізмів приділяється велика увага. Так, А. Емірова [6, с. 62] називає іменники "голова", "серце" та ін. "центрами, що організують цілі гнізда стійких зворотів". На думку Ю. Аваліані, іменник "голова" служить центром типологічного ряду із загальним значенням оцінки розумових та моральних якостей особистості і є стрижневим словом для цього типологічного ряду. При утворенні подібних рядів ці стрижневі слова "по суті виконують роль індикаторів, котрі класифікують знаки певної теми, яка розкривається в подібних рядах" [1, с. 2-3].

Метою нашої розвідки є окреслення особливостей семантичного співвідношення розуму та почуттів, представлених у польській фразеології, і визначення, з якими асоціативними образами пов'язується в носіїв польської мови позначення розумових здібностей людини та почуттів особистості.

Актуальність дослідження зумовлена тим, що фразеологічні звороти певної тематичної групи завдяки своєму широкому використанню як у розмовній, та в

© Дем'яненко Н., 2010