

KORPUS A KORPUSOVÁ LINGVISTIKA

FRANTIŠEK ČERMÁK

KAROLINUM

Korpus a korpusová lingvistika

František Čermák

Recenzenti:

prof. PhDr. Eva Hajičová, DrSc.

PhDr. Jitka Šonková, CSc.

PhDr. Věra Schmiedtová

Vydala Univerzita Karlova

Nakladatelství Karolinum

Redakce Lenka Ščerbaničová

Grafická úprava Jan Šerých

Sazba DTP Nakladatelství Karolinum

Vydání první

© Univerzita Karlova, 2017

© František Čermák, 2017

ISBN 978-80-246-3710-5

ISBN 978-80-246-3776-1 (online : pdf)



Univerzita Karlova
Nakladatelství Karolinum 2017

www.karolinum.cz
ebooks@karolinum.cz

Slovo úvodem

Text knihy zachycuje vznik a rozvoj korpusů a korpusové lingvistiky nutně k určitému datu, vývoj jde však rychle dál. Vyjadřovalo se k němu více lidí, můj hlavní dík však patří jejím recenzentům, kterými byli prof. PhDr. Eva Hajičová, DrSc. (MFF UK), PhDr. Jitka Šonková, CSc. (University of Iowa) a PhDr. Věra Schmiedtová (FF UK).

František Čermák

OBSAH

I. ČÁST: TEORIE ----9

1	ÚVOD: KORPUS A KORPUSOVÁ LINGVISTIKA ---- 10
2	DATA A INFORMACE ---- 12
2.1	Elektronický text ---- 13
2.1.1	Obecná povaha, vlastnosti a zpracování textů ---- 13
2.2	Korpusová data ---- 16
2.2.1	Povaha a zdroje elektronických dat ---- 17
2.2.2	Recepce a produkce textu ---- 19
2.2.3	Reprezentativnost korpusových dat jako zrcadlo reality ---- 21
2.3	Korpusový kontext ---- 25
2.4	Korpus a empirie ---- 27
2.5	Technické a další aspekty přípravy korpusu ---- 28
2.5.1	Příprava dat pro korpus ---- 29
2.5.2	Dotazy a dotazovací jazyk ---- 34
3	KORPUS A KORPUSY ---- 37
3.1	Korpus, jeho povaha a vytváření ---- 37
3.1.1	Korpus a jazyk ---- 37
3.1.2	Tvorba korpusu a jeho možnosti ---- 39
3.1.3	Hledání v korpusu a jeho analýza ---- 41
3.1.3.1	Možnosti hledání a analýzy ---- 41
3.1.3.2	Souhrnný pohled na úzus slova v kontextu: Konkordance ---- 51
3.1.3.3	Kombinace textové i systémové: kolokace a koligace ---- 55
3.1.3.4	Postup od formy k významu a funkci. Typický a zvláštní ---- 62
3.1.4	Hlavní aspekty a charakteristiky korpusu ---- 67
3.2	Český národní korpus (ČNK) ---- 69
3.3	Typy korpusů ---- 74
3.3.1	Synchronní korpus ---- 76
3.3.2	Mluvený korpus ---- 78
3.3.3	Diachronní korpus ---- 80
3.3.4	InterCorp ---- 82
3.3.5	Webové korpusy ---- 83
3.4	Hlavní světové korpusy ---- 84

4	KORPUSOVÁ LINGVISTIKA ---- 90
4.1	Korpusová lingvistika jako disciplína ---- 90
4.2	Korpusová metodologie ---- 92
4.3	Korpusová statistika ---- 100
4.3.1	Základní pojmy a operace ---- 101
4.3.2	Frekvence ---- 102
4.3.3	Zipfův zákon ---- 106
5	KORPUS VE VÝZKUMU A APLIKACI ---- 109
5.1	Výzkum spojený s budováním korpusu ---- 110
5.2	Lingvistický výzkum ---- 111
5.2.1	Výzkum obecně ---- 111
5.2.2	Variabilita jazyka. Mluvená povaha jazyka ---- 112
5.2.3	Lexikon a frazeologie. Kolokace ---- 114
5.2.4	Morfologie a gramatika. pedagogické aplikace ---- 117
5.3	Lingvistický kontrastivní výzkum více jazyků (InterCorp) ---- 118
6	ZÁVĚR: PERSPEKTIVY KORPUSOVÉ LINGVISTIKY ---- 120

II. ČÁST: SONDY, STUDIE, ANALÝZY ---- 123

A	Korpusy včera, dnes a zítra (2011) ---- 125
B	Some of Current Problems of Corpus and Computational Linguistics or Fifteen Commandments and General Truths (2007) ---- 145
C	Spoken Corpora Design: Their Constitutive Parametres (2009) ---- 156
D	InterCorp: jeho povaha a možnosti (2014) ---- 165
E	Lexical Collocability: The Case of Verbs and Adverbs (2010) ---- 182
F	Abstract Nouns Collocations: Their Nature in a Parallel English-Czech Corpus (2005) ---- 196
G	Statistical Methods for Searching Idioms in Text Corpora (2006) ---- 207
H	Povaha a úzus interjekcí: případ češtiny (2007) ---- 216
I	Konference: Staré známé nebo neznámé slovo? (2001) ---- 225

III. ČÁST: PŘÍLOHY ---- 229

Příloha A	Jazyková analýza výběrové konkordance (linout se) ---- 231
Příloha B	Frekvence slov ve Frekvenčním slovníku (prvních 100 lexémů) ---- 234
Příloha C	Jeden pohled na českou morfologii (Procentuální rozložení rodu, čísla a pádu u substantiv) ---- 237
Příloha D	Román jako korpus: Zbabělci (Josef Škvorecký) ---- 239

KORPUSOVÁ BIBLIOGRAFIE ---- 247

GLOSÁŘ ZÁKLADNÍCH POJMŮ A TERMÍNŮ ---- 261

I. ČÁST: TEORIE

1 ÚVOD: KORPUS A KORPUSOVÁ LINGVISTIKA

Tato kniha je stručným shrnutím a přehledem základních pojmů, zásad a metod **korpusové lingvistiky**, tj. nauky o korpusu a práci s ním, i nového stejnojmenného oboru. Už podle svého jména se tu staví nutně na **korpusu**, tj. velkém elektronickém, systematicky budovaném a organizovaném úhrnu textů, který je *conditio sine qua non* pro jeho disciplínu, tj. korpusovou lingvistiku. Kniha zároveň přináší i malou ilustraci možností výzkumu s modelovými výsledky.

Korpus se objevil v lingvistice, spolu s nástupem počítačů, poměrně nedávno a výzkum jazyka v ní v důsledku existence korpusu zcela převrátil poměry a etablované postupy, především z hlediska rozsahu informací a možností, které data i počítače nabízejí. Dalo by se tu, jakkoliv se to běžně nedělá, mluvit o korpusové revoluci v lingvistice, která jde mj. ruku v ruce s rozvojem počítačů, aspoň na svém počátku.

Jestliže lingvista tradičně donedávna strádal zásadním nedostatkem dat, a tedy, metaforicky řečeno, jistou informační podvýživou, pak se s nástupem korpusů situace radikálně změnila. Tam, kde v celé dlouhé minulosti oboru lingvista pracně, dlouho a nedokonale shromažďoval svá manuální excerpta v podobě lístečků a výpisků, než je mohl začít třídit a dávat do své kartotéky (a později i většího archivu), kterými se obv. pak sekvenčně probíral, tam nastoupil počítač a jeho možnosti. Na místě někdejších lístečků a výpisků dnes korpus poskytuje doslova záplavu informací a dat, ve kterých je sice často obtížné se vyznat přímo a snadno, ve kterých se ale dá už hledat více způsoby, vždy pomocí počítače. To vedlo a vede mimo jiné k nutnosti nalezení (a dalšímu hledání) potřebných technik a metodologie, jak k tomuto množství dat smysluplně přistupovat.

Korpusové poznání přináší v mnoha ohledech překvapivá zjištění. Tato kniha výběrově přibližuje možnosti mnohostranného poznání všeho základního, co se událo za cca dvacet a více let v této nové disciplíně, **korpusové lingvistice**, jejíž

výsledky a aplikace začínají i u nás už otřásat tradiční důvěrou ve spolehlivost stávajících lingvistických prací, starých slovníků, mluvnic a dalších příruček. Mimo jiné se korpusová lingvistika postupně stala už i novým lingvistickým oborem studovaným na univerzitách.

Kniha postupuje od širšího rámce, potřeb moderní lingvistiky, povahy korpusu a jeho možností až k některým základním a konkrétním autorovým výsledkům z jeho studia a výzkumu v daném oboru. Svým systematickým důrazem na vyčerpávající přístup k datům a jejich popisu, který v lingvistice přináší a důsledně uplatňuje poprvé korpusová lingvistika, se tato kniha snaží korpusová data a přístupy k nim aspoň základním způsobem charakterizovat a narýsovat tak zde zároveň i půdorys této rychle se rozvíjející disciplíny. Stejně tak se snaží ukázat na průnik i návaznosti na lingvistické disciplíny tradiční a ovšem přitom i naznačit otevřené otázky a problémy. Jakkoliv se hlavní zkušenost opírá o práci s budováním a analýzou **Českého národního korpusu** (viz blíže na adrese *korpus.cz*), v pozadí stojí i zkušenost s většinou korpusů větších a řadou menších, především evropských, a kontakty s týmy, které je budují, která se prezentuje výběrově také.

Tato stručná oborová příručka a malý sborník zároveň se pokoušejí komplexně zachytit tedy jak povahu korpusových dat, tak způsoby jejich zpracování, možnosti práce s nimi i všechny hlavní souvislosti a návazné aspekty tak, aby podala stručný komplexní a ucelený, avšak zároveň i základní obraz **oboru** korpusová lingvistika (**I. ČÁST: TEORIE**).

Doprovázejí ji a jen volně na ni navazují reprinty vybraných autorových studií a analýz představující některé výsledky analýzy korpusu (**II. ČÁST: SONDY, STUDIE, ANALÝZY**), které na konci doplňuje ilustrativním způsobem několik konkrétně orientovaných dodatků (**III. ČÁST: PŘÍLOHY**).

Svým pojetím je kniha určena jak studentům a lingvistům či odborníkům z jiných oborů, především humanitních, tak i širší zainteresované veřejnosti. Kniha je vybavena dvojí bibliografií, obecnou, vztahující se k celé problematice disciplíny, a specifickou, vztaženou přímo k tématům v podobě stručných odkazů za jednotlivými kapitolami. Bibliografie tak do značné míry vyvažuje i specifikuje stručnost formulací a popisů zde obsažených a orientuje uživatele na další možnosti studia. Příspěvky v části II mají původní bibliografii vlastní.

ZÁKLADNÍ LITERATURA

Čermák 1995a, 1998, 2001b; McEnery – Wilson 1996; McEnery – Hardie 2012; Wynne 2005.

Speciální a další literaturu viz v Korpusové bibliografii

2 DATA A INFORMACE

Informace se chápe různě a ve více smyslech, intuitivně však především jako nějaký údaj, zjištění o něčem, resp. už i přímo zobecněné poznání či vědění. Informace se odborně někdy vymezuje jako to (tedy zjištění, poznatek), co relevantního lze získat z podkladových dat (údajů) pro analýzu jevu či problému, a to obvykle v nějakém rámci, pro lingvistiku pak z dat textových, a to zvláště v rámci a kontextu sociálním. Ten se v případě jazyka chápe především jako rámec sémiotický a je řízený pragmatickými pravidly.

Současný informační věk a jeho povahu obecně i specificky reflektuje i korpus: obojí, tj. náš věk i korpus sám, nabízí nebývalou záplavu informace. Její množství, které je zásadně vítané, však i hrozí zahlcením uživatele, a představuje proto také relativně nový problém. Pravda je, že náš informační věk se už dobře neobejde bez počítačů a globálního webu, na tyto nástroje a **hromadné** zdroje informace spoléháme dnes zcela samozřejmě. Této novodobé pravdě ovšem předcházela a předchází pravda starší a stále zcela základní, totiž že **valná většina informace** (významová, formální i pragmatická) **je kódovaná a přenášena (přirozeným) jazykem** a je jen málo oborů, které se při výměně informací bez jazyka obejdou (srov. část matematiky, popř. i malbu apod.); na to se často zapomíná.

Tyto informace, často po určitém zpracování a formalizaci, se obvykle předávají dál, dalším uživatelům, zájemcům. Pak se role toho, kdo informaci objevil, získal (zvl. odborník, lingvista), mění a stává se z něj ten, kdo ji předává, sděluje v nějaké podobě jiným lidem. Tyto informace, tj. poznatky zvláště o aktuální povaze a stavu toho, co nebo kdo nás zajímá, se dnes výrazně sdělují elektronicky, a to často **hromadně**.

Stále se však dosud a hlavně sdělují i po staru také **individuálně**, mj. ústním a soukromým kontaktem mezi lidmi, osobním pozorováním, zkoumáním, poznáním něčeho, popř. i zčásti zobecněnou a zprostředkovanou školní zkušeností, a tedy vlastní deduktivní či induktivní činností. Taková informace, např. článek, kde lingvista zformuluje své poznání a zjištění, či ho odpřednáší studentům, se tak dostává dál. Je však třeba lišit, zda takové poznání autor považuje, po ověření jeho platnosti, tj. zda skutečně odpovídá datům a obvyklým přijímaným normám, popř. i stupni poznání, za zobecnitelné, za obecně přijatelné a považuje ho tedy za poznání a informaci **objektivní**. Proti poznání objektivnímu jde však často i o poznání, zkušenost i informaci **subjektivní**, vázanou na jednotlivce apod., kde míra zobecnitelnosti nebývá jasná.

Oba zdroje poznání a informace z něj odvozené, zdroj individuální i hromadný, jsou ovšem komplementární a jeden nevylučuje druhý. Nicméně individuální a často subjektivní zkušenost a informace nabývaná jedincem obvykle nemívá

možnost **potvrzení** své správnosti ani obecnosti, a tedy zjištěné **platnosti** při opakovém výskytu, tj. ověřované opakovaně, a použitelné tedy i ostatními. Taková povaha informace se naopak očekává, ba vyžaduje u informací veřejných, obecných, a zvláště vědeckých.

Data, informační jednotky a základní stavební kameny poznání, jsou podle oboru a druhu vnímání a zapojeného smyslu pochopitelně různá, (téměř) všechna jsou však převoditelná do slov, a tedy jazyka, psaného a mluveného, který nám zprostředkuje zrak a sluch. A až tady se zúročuje možnost, zprostředkovaná komputerově, dostávat se k zobecnění, popř. zobecnitelné, a tedy i zpravidla **objektivnější a opakovatelné**, resp. opakované informaci, tj. takové, kterou lze nalézt až ve větších informačních jazykových souborech, resp. textech, tj. kterou v nich lze ověřit a kterou individuální a jednorázová zkušenost a poznání z vlastního výzkumu v zásadě nenabízí.

Toto **zobecnění** jednotlivého *faktu*, tj. poznání, ať je induktivní či deduktivní, a jeho *platnosti* je ovšem možné jen v **kontextu** a na základě či proti pozadí široké znalosti dané vzděláním obecným či oborovým; izolovaný fakt nám naproti tomu neřekne (téměř) nic.

Data a informace v nich však ještě nevedou k tomu, co je zvláště v případě seriózního studia na jeho konci to nejdůležitější, to jest k **věděni**, popř. znalosti, jak informace používat a využívat. To je však už věcí integrace poznatků do platného kontextu a jejich nazírání v něm, ale i studia a učení, opřené o opakování a postaveného kvůli souvislostem do vhodného širšího rámce, zvláště celého jazyka i disciplíny.

2.1 ELEKTRONICKÝ TEXT

2.1.1 OBECNÁ POVAHA, VLASTNOSTI A ZPRACOVÁNÍ TEXTŮ

Základní a nejčastější podoba informace je obsažená v *textu*, dnes většinou už elektronickém, tj. v textu vytvořeném pomocí počítače (a jím i dále zpracovávaném). Každý text je smysluplný řetězec znaků, pro lingvistiku a uživatele jazyka řetězec hlavně slov, který má lineární povahu. Psaný text má ale i povahu lineárně prostorovou, kde platí dimenze *vlevo-vpravo* (podle druhu písma i jiná), nebo lineárně časovou, kde platí v podstatě dimenze *napřed-potom* (u mluveného jazyka se jeho povaha lineárně časová přepisem mění na lineárně prostorovou taky). Paralelní jiná informace (vedle základní psané), zvláště intonace u mluvené komunikace (suprasegmentální rysy), se v písmu nezachycuje, a jen mluvený text je tudíž vlastně vícelineární, skládá se z více souběžných informačních linií (více viz 2.3); mluvený text je však také obecně primární a historicky starší.

Elektronický text je pro potřeby korpusové lingvistiky text, který je počítačově čitelný a dále zpracovatelný. V lingvistice se za text považuje libovolný souvislý text, získaný obvykle jen z autentických zdrojů, v praxi konkrétně (smysluplný) text psaný nebo mluvený. Míra a způsob zpracování a přípravy elektronického textu závisí na cíli, technických podmínkách a potřebách. Zpracování textu psaného a mluveného se liší. Podoby psaného textu se však liší také.

Obecně je při zpracování textů třeba lišit tři základní pojmy, jejichž míra obecnosti se postupně zužuje, a to nepřímou úměrně k nárůstu jejich specifčnosti: typ dat, kód dat a formát (obv. souboru). **Typ dat** je pojem nejširší a patří sem jen okrajově. Tvoří ho proměnlivý *způsob ukládání dat různého druhu*, dat vizuálních, obrázků (např. *jpeg*, *gif*) či akustických (např. *wav*, *wma* či kompresní *mp3*). Pro korpusy jsou však zásadní druhé dva pojmy, kód (kódování) dat a jejich formát (obv. formát souboru).

Texty se zachycují, resp. znázorňují v různých **kódech**, které v elektronické oblasti označují *způsoby jejich převodu (a informace v nich) do způsobu jiného*; patří sem např. morseovka (písmena převáděná na tečky a čárky), (dříve) dírky a jejich uspořádání na děrném štítku, a dnes konečně a hlavně čísla (na která se převádějí písmena). V této poslední oblasti (užívající číselného kódu) jsou dnes běžné dva typy kódů a kódování, ASCII a Unicode. ASCII, které prodělalo více proměn a postupné rozšiřování, dnes čítá jen pro češtinu až pět různých verzí, a protože obecně stále není plně použitelný pro všechny jazyky, přechází se proto, resp. přešlo se do univerzálního *Unicode* (viz víc v 2.5.1 a 2.1.1).

Naproti tomu **formát** je *způsob technického záznamu a zpracování informace v určitém kódu technických dat*, který je značně proměnlivý. Formátů může být pro uložení téže informace víc než jeden (vedle netextových formátů, viz výše např. *mpeg*), specificky mj. *doc*, *T602*, *ODF pro Linux*, *html*, či obecný *XML*, i podle druhů operačního systému, pro Windows, Unix či Mac OS). V tomto smyslu tedy formát (resp. způsob kódování znaků, viz dále) je *takové uspořádání dat, které slouží k jejich znázorňování, ukládání a znovuužívání*. Známý je mj. formát *pdf* (*portable document format*) užívaný pro různé typy textů jako snadno převoditelný, určený zvláště pro tisk (s informací o charakteristikách takového textu, ale i obrázky), který je známý z běžné práce s elektronickými texty.

Elektronické texty, které dnes vstupují jako jeho stavební kameny do korpusu, se tedy vyskytují v různých podobách, *formátech*, a pro potřeby počítače se musejí konvertovat do formátu jednotného, tj. sjednotit. Vedle formátů spojených s operačním systémem (*Windows*, *Unix/Linux*, dříve *DOS*) se texty vyskytují ve formátu, který jim dávají různé textové editory a procesory (jako *Word*, *Writer* v *LibreOffice* či *TextEdit* v *Mac OS X* aj.), existují ve formátech vznikajících i ve webovém prostředí (zvl. *html*, *Hypertext Markup Language*) a ty je třeba podrobovat konverzi (viz 2.5.1).

V praxi se po konverzi a převodu ze specifického formátu výsledný formát označuje obvykle jako *txt* (prostý text). Zachycuje ho rozšířený kód ASCII (pův. *American Standard Code for Information Interchange*), obvykle už v univerzálním kódování Unicode (viz 2.5.1), specificky v UTF-8. Tak se tvoří základ pro další korpusové zpracování a budování vlastního korpusu (srov. dál ještě následnou úpravu textu jazykem XML, zahrnujícím i další zpracování, 2.5.1).

Vybraný korpusový text, získaný různými konverzemi textů z textových procesorů užívajících různých formátů, je ve výsledku zásadně text *autentický*, a má tu podobu, kterou mu dal autor, resp. vydavatel. Je tedy nepředstavitelné, že korpusový text, který by měl ideálně zachovávat záměr autora, by se měl opravovat, dál editovat, pospisovňovat či jinak pravopisně, nebo dokonce cenzorně, a tedy krajně sporně „vylepšovat“ (což je věcí konkrétní a vždy problematické jazykové politiky a inženýrství). Ochranu autora a autentičnost jeho textu je třeba dodržet za každou cenu (jakkoliv i autorská formální úprava textu může být nevyrovnaná a naznačovat problémy tvůrce textu s jazykem). Předpokládá se tedy, že korpusový text zachovává i různé individuální překlery a omyly autora, což pro masové korpusové hledání nevadí. Takové, v tradičním pohledu „spisovníků“ a tradicionalistů, nenáležitosti bývají totiž řádké a neovlivňují zásadně celkový výsledek hledání v korpusu (viz 2.5.1). Korpusy mají tedy mj. i „konzervační“, archivní roli tím, že uchovávají (relativně) trvale texty v neměnné podobě (jakkoliv ale přitom z hlediska typografického už není jejich původní podoba zpětně rekonstruovatelná). Při přípravě textů pro korpus se však vyskytují i (často hromadné) chyby textu vzniklé technickým zpracováním či sazbou textu (a vzniklé tedy bez, anebo proti přání autora), a ty je naopak třeba zjistit a řadou procedur opravit, jakkoliv rozlišení obou typů nebývá snadné.

Protože zásadní podobou korpusových dat je stále text psaný, spadají sem nutně i *elektronické přepisy* mluvených textů (viz 3.3.2), které mají také textovou podobu. Pro zachování jednotnosti s texty psanými mají být s nimi ideálně srovnatelné, avšak hledání jejich optimální formy (v zásadě formy transkripce) při zachování jejich věrné původní podoby je zatím otevřená otázka (zvl. s ohledem např. na nespisovné či individuální podoby slov). Takové texty se můžou dodatečně ještě vybavovat i fonetickým či prozodickým přepisem (zachycujícím jejich suprasegmentální rysy).

Proti dřívější koncepci budovat korpus z homogenních (náhodných) **vzorků** textů (ty vyhovují zvl. statistickým potřebám někdy lépe) se dnes dává, mj. i vzhledem k dostupnosti elektronických textů, přednost tomu, do korpusu začleňovat **texty celé**, které umožňují optimální zkoumání kontextu, stejně tak ale i povahu odlišných částí téhož textu (jako je specifický začátek apod., při textové analýze). Problém, kdy se může přístup založený na vzorcích textů jevit výhodnější, se např. projevuje tehdy, když jde o malou textovou kategorii či (pod)obor, jehož naplnění

více různými malými vzorky je lepší a vhodnější než jediným a úplným textem větším. Ani velké množství náhodných vzorků však nemusí zachytit některé zvláštnosti textu.

2.2 KORPUSOVÁ DATA

Protože většina korpusů je veřejných a zájemcům nekomerčně přístupných, bývá získávání elektronického textu (*akvizice*) od poskytovatelů spojeno s nemalými problémy autorských práv, **copyrightu**, který je podmíněn souhlasem různých agentur (zastupujících autory, tj. vlastníky textů) udělujících souhlas k využití textů na základě zvláštních smluv. Někdy je dokonce jejich využívání korpusu i zpoplatněno (stranou zůstávají soukromé či podnikové korpusy, veřejnosti nedostupné, vyvíjené zvl. v zahraničí například některými nakladatelstvími, kde může být autorský souhlas delegován na tato nakladatelství apod.). Ať už jde přímo o nakladatelství, vydavatelství či texty zprostředkující organizace a agentury, všechny jsou vždy odpovědné autorům za to, že se dodržuje autorský zákon v případě textů, se kterými pracují, i za jejich ochranu, tj. zajišťuje dodržování licenčních podmínek, za nichž byly poskytnuty; neautorizované šíření textů je právně postižitelné. Korpusová centra přebírají tyto texty sice z různých zdrojů, ale vždy za podmínky respektování domluvených podmínek, především ve smyslu ochrany před dalším šířením, a to na základě právně závazných smluv.

Případně širší, tj. zvláště technické a tedy nelingvistické využití autorských textů, které tvoří korpus, je věcí dohody takového zájemce a poskytovatele textů (zastupujícího autora textů). Avšak i standardní **využití** korpusových textů, které jsou pro korpusové užívání provozovateli korpusu, tedy pro dané využití, jen svěřené (v českém případě je to zvl. Ústav Českého národního korpusu Filozofické fakulty Univerzity Karlovy), je omezené na dohodnutou limitovanou délku citace (zvl. v podobě rozsahu, pozic či vět ap.). Takové užívání je v zásadě otevřené a u konkrétního lingvisty apod. vázané na respektování podmínek užívání. Pro odborný technický výzkum textů se po dohodě může zájemcům nabídnout i *podoba textů s namátkově přeházenými větami* či jinými úseky (*reshuffled texts*), kde už nejde přímo o autorský text v původní podobě, a práva autorů se tak v zásadě nenarušují.

Ekvivalentem, obdobou copyrightu je u mluvených korpusů *dohoda* s nahrávanými osobami o podmínkách zveřejnění nahrávek, kde právně namísto autorského zákona začíná platit občanský zákoník. Nemalý problém, narážející na nestejnou legislativu autorských práv v různých zemích, se musí řešit při výstavbě paralelních (vícejazyčných) korpusů. Více o vlastnictví textů viz též v 2.5.

2.2.1 POVAHA A ZDROJE ELEKTRONICKÝCH DAT

Elektronická data se zpracovávají z více druhů zdrojů z dodaných, do počítače vnesených textových dat (viz 2.5.1). Dnes je u psaných textů *synchronních* nejčastější, nejsnadnější a nejužitečnější ta podoba elektronická, která se získává, často přes nějakého zprostředkovatele (nakladatelství, agenturu ap.), a je už produkovaná přímo v elektronické podobě (většina autorů dnes sama píše na počítači v nějakém textovém editoru).

Naproti tomu se *starší* texty musejí zpravidla ručně a pracně skenovat (programy OCR, optical character recognition), zatímco *mluvené* texty se nejčastěji do počítače napřed a neméně pracně přepisují z akustického záznamu, resp. nahrávky a až pak dál zpracovávají jako jiné elektronické texty.

Korpusová data, prostá či (většinou) různě anotacemi obohacená a technicky zpracovaná, jsou, co do zdroje informace, zhruba řečeno moderním elektronickým ekvivalentem starých a rozsahem omezených manuálních excerpt z různých (většinou psaných) zdrojů (v podobě kartotéčních lístků). Z nich se získávala a někdy ještě i získává prakticky veškerá informace o jazyce, tj. generalizací z více podobných či stejných výskytů (lecko si dokonce i dnes dělá malé kartotéky dosud sám). Korpusová data jsou, jinými slovy a zjednodušeně řečeno, rozsáhlé soubory elektronických textů, které jsou následně uspořádány a propojeny do podoby korpusu. Tato data jsou na rozdíl od soukromých osobních, omezených, a tedy nutně subjektivních excerpt jednotlivce zásadně *objektivní*, protože ve velkém a zvláště reprezentativním korpusu odrážejí v širokém průřezu nejrozumnější a typologizované druhy textů, a tedy jak většinový a *typický* úzus jazykových forem, tak ale i úzus menšinový (daný nižší frekvencí jejich výskytu), který zrcadlí jejich různé a různé časté zastoupení v korpusu. **Reprezentativní podoba** dat v korpusu se obvykle považuje v nespécifickém a globálním výzkumu jazyka usilujícím o celkový a vyvážený obraz úzu za optimální. Zprůměrování získaných výsledků naznačuje pak i typický úzus.

V kontrastu k tomuto většinovému přístupu korpusové lingvistiky usilujícímu o objektivnost informací a podkladových dat stojí dosud některé přístupy např. N. Chomského a jeho stoupenců. Ten objektivitou informace, kterou korpusy nabízejí, v zásadě pohrdá a upřednostňuje své vlastní izolované a často i spekulativní či konstruované příklady úzu, na kterých zakládá své argumenty a proměnlivé teorie. Vyvozovat zobecnění z jednotlivých a izolovaných příkladů, které takový přístup nabízí, je však krajně problematické a stojí v protikladu k cílům korpusové lingvistiky analyzovat text důsledně celý a systematickým způsobem, bez vynechávek.

Každý korpus je dnes založený na souborech elektronických **dat psaných**, připravených pro tisk nebo internetovou komunikaci různého druhu (jiná psaná data

je třeba skenovat, viz už výše) či **dat mluvených**. Psané texty se dnes získávají většinou (i když ne vždy) od vydavatelů tištěných textů na základě dohody s nimi a kvůli jejich velikému rozsahu jsou také automaticky zpracovávány, i když mívají různé formáty; v řadě případů je třeba při získávání dat pro korpus autorská práva tvůrců či majitelů textů chránit (copyright, viz i výše).

Náročnější je proces získávání **mluvených dat** pro mluvené korpusy (viz 3.3.2); představuje ho především nahrávání zvláště spontánních promluv. Takové nahrávky se následně přepisují do grafické podoby tak, aby mohly případně existovat paralelně s akustickou podobou. Jiné, **vícemodální korpusy** (zahrnující i obrazy, digitální, někdy i filmové, ve snaze zachytit dynamiku zachycované situace promluvy) pro náročnost jejich získání i zpracování dosud ve významnějším rozsahu neexistují, a jsou proto stále víceméně zbožným přáním. Problém je nejen v metodologii záznamu (propojení obrazu, zvuku a písma), způsobu hledání v nich, ale i takových věcech, jako je způsob anotace gest, posunků, výrazů tváře, typů situace, nálady či postojů mluvčích ap., tedy všech **mimotextových faktorů**, které se spolupodílejí na významu a funkci; žádná shoda napříč jazyky a týmy (pokud se věci vůbec zabývají) není. O takovém korpusu, srovnatelném svou mírou informativnosti a objektivnosti s korpusem psaným, nemůže být proto zatím ani řeč, není tu dostatečné množství slov ani kontextů. Přestože existuje pár pokusných vícemodálních korpusů, které implicitně ovšem v praxi zatím vlastně znamenají *pouze dva zaznamenané mody* podle lidských smyslů uplatňujících se při vnímání nejčastěji, tj. zraku a sluchu (např. film představuje jejich spojení), lze se pro budoucnost ptát na možnost zapojení i smyslů dalších, zvl. čichu a hmatu. Nejde tedy tolik o mody (a název multimodální není vhodný), jako o druhy percepce (a mělo by se event. mluvit o korpusech multiperceptuálních). Nikdo si však otázku, kolik módů komunikace zachycovat a jakým způsobem, vlastně dosud neklade, zřejmě proto, že neví, jak by na ni odpověděl.

Od korpusových dat, korpusů se liší **archivy**, původně manuální, dnes však už většinou také elektronické, které bývají veřejné (jako *Oxford Text Archive*), nebo vázané na lexikografická, popř. jiná centra (jako *Archiv ÚJČ ČSAV*), která svá data shromažďovala už dříve, a to ruční excerpcí (na kartotéční lístky); jejich význam už dnešní korpusy v podstatě odsunuly do pozadí.

Archivy, konzervující na rozdíl od korpusů jakákoliv cenná elektronická data, ať už diachronní nebo mluvená, jsou však spíše výjimkou. V druhém případě, tj. u mluvených dat, jde však mnohem víc o výjimku, protože automaticky spolehlivý převod akustických dat (z existujících archivů) do psané podoby dosud k dispozici není. V prvním případě, i když je jejich povaha také manuální, archivy se stále udržují a někdy i rostou dál (jako *archiv Oxfordského slovníku* sestavovaný přes 150 let na základě jeho *Reading programme* s pomocí původně dobrovolníků výběrově excerpujících dodané knihy). Starší data se však i zde musejí převádět

do elektronické podoby (viz 3.3.3) většinou náročným ručním skenováním a procházejí pak dalšími procesy.

Významným typem aktivity, obecně sdíleným v mnoha zemích, je ve snaze o zachování národního kulturního dědictví úsilí o naskenování národní literatury minulé i nové v knihovnách, popř. i dalších textů, které jsou (jednoduchým způsobem) pak na webu veřejně a obecně zpřístupňované. Tak postupuje pro českou oblast např. Národní knihovna. Často se však takový archiv skládá, pro snadnější způsob jeho vytvoření (protože kvalita automatické digitalizace starších textů není uspokojivá), jen z „obrázků“, tj. stránek textu naskenovaných jako celek (a tedy obrázek), ve kterých slova obvyklým způsobem hledat nejde, a pro korpusové využití taková data nepostačují.

K neznámějším obecným a velkým mezinárodním zdrojům textů patří výsledky činnosti několika organizací, především *Project Gutenberg* (Open Language Archives Community), *Elsevier.org* či *OTA* (Oxford Text Archive) či u nás *Národní knihovna* (srov. <http://www.nkp.cz/digitalni-knihovna>, zvl. skrze aplikaci *Kramerius*, <http://kramerius.nkp.cz/kramerius/>). Některé další jsou komerční jako *ELRA*, *ELDA* či americká *LDC* (adresy viz v bibliografii).

2.2.2 RECEPCE A PRODUKCE TEXTU

V autorství psaných textů se projevuje jasná disproporce: je mnohem víc pasivních uživatelů (čtenářů, konzumentů) než těch aktivních (autorů, tvůrců), tj. těch, co texty vytvářejí, píšou. Ideální stará představa bývala, že korpus by měl být ve vyvážené proporcii složený z textů odpovídajících (širokému publiku a tedy) recepci i produkci, tj. mělo by být zohledněno i množství autorů textů; tak se měl specificky odrážet i jejich soukromý jazyk (ne nutně však veřejný). V korpusové lingvistice (i jinde) nejde tedy pouze o uživatele v první řadě, nýbrž i o autory, tvůrce (korpusových) textů. Tento rozdíl se nazývá **recepce textu** a stojí proti **produkcí textu**; je to specifický problém především textů psaných. U mluveného jazyka a mluvených textů tento rozdíl do té míry nevystává, všichni mluví i poslouchají zároveň, jakkoliv ani zde ne ve zcela stejné míře (ale to už je věc situace a psychologicky i povahy mluvčích; je tu i nejednoznačný vliv rozhlasu). Původní představa, že by u psaného jazyka korpus měl nějak odrážet proporce produkce a recepce, byla zkusmo uplatněná v zásadě jen u dánského korpusu, kvůli chybějícímu výzkumu a dostupnosti reprezentativního výběru textů podle produkce se však dnes (zatím) neuplatňuje.

Pokud chtějí mít reprezentativní podobu, budují se korpusy v praxi tedy především podle kritéria **recepce**, tj. rozsahu a složení čtenářstva (psaných) textů a povahy jimi čtených textů. To je různě velké, největší u novin, nejmenších zřejmě u odborných tisků či tisků administrativních a podle obsahových hledisek je i proměnlivé. Tyto

texty se dají relativně snadno získat, jejich proporce promítající se do reprezentativního složení korpusu jsou však stále věcí diskuze. Proti textům z hlediska recepce stojí texty získávané podle **produkce**, tj. texty ideálně pocházející od (co nejvíce) různých autorů a posuzované i podle množství těchto autorů; tento ohled tvůrci korpusů při výběru textů sledují jen málo a takové texty bývají v menšině a je obtížné je získávat ve vyvážených proporcích.

Obecně jde však také a z jistého hlediska především o *formování, ovlivňování čtenáře* a jeho jazyka skrze texty. Při bližším pohledu je zřejmé, že na jedné straně je relativně velmi málo lidí (autorů, spisovatelů, novinářů), kteří svou tvorbou zásadně a *aktivně* ovlivňují svou produkcí (psaný) jazyk a jeho konzumaci čtenáři (jakkoliv významným a blíže nezkoumaným způsobem). Proti tomu však stojí většinový úzus platný pro obrovskou většinu čtenářů na straně druhé, kteří jsou v roli pasivních čtenářů textů, a tedy i uživatelů korpusu. Platí to v podstatě i obecněji, srov. např. čtenáře bulvárních novin (s nejvyššími náklady) vytvářených malou skupinou novinářů pro obrovskou masu čtenářů (přitom erudice a jazyková kompetence těchto novinářů nemusejí být nejlepší).

Určité omezené možnosti vyvážení této jinak přirozené disproporce (zvl. v dialogu) představují *speciální korpusy*, jako je např. korpus korespondence (viz 3.3), zaměřené jen na jeden speciální druh textu, jehož uživatelé jsou odpovídajícím způsobem specializovaní také; zpravidla to jsou lingvisti a literáti.

Moderní korpusy, zvláště ve své publicistické složce, však nemusejí být zcela jednojazyčné, i když jednojazyčnost je předpoklad samozřejmý. V důsledku reklamy (někdy i rozsáhlejší), popř. **cizojazyčných** citátů (kombinace slov i celé věty) se v novinách apod. objevují i menší texty cizojazyčné, zvl. z jazyků velkých a známých, tj. angličtiny, němčiny, latiny apod. (někdy celé odstavce). Tyto nečeské texty do textu českého bývají ovšem vloženy (zvl. z různých důvodů nakladatelem či vydavatelem) a text jejich dodatečnou přítomností pak ztrácí svou autentičnost. Je pochopitelné, že pro národní jednojazyčný korpus takovéto texty nejsou žádoucí a musejí se odstraňovat. Rysem této cizojazyčné složky je malá či žádná opakovanost a (i pro český korpus) jistá bezkontextovost (resp. slabá, popř. odlišná závislost cizích slov na kontextu jiného jazyka, např. němčiny na českém kontextu). To se ovšem netýká *citátových cizojazyčných slov a výrazů*, ať jde o módní anglicismy či klasické latinské či řecké výrazy, srov. *à propos, acquis communautaire, last but not least* aj., které se do českého úzu svou vysokou frekvencí už zařazují a užívají se i v širším českém jazykovém kontextu. Okrajovým případem ovšem je český text, který autor cizojazyčným textem vybavil sám; ty je pak třeba, v zájmu zachování autentičnosti, v korpusu ponechat.

2.2.3 REPREZENTATIVNOST KORPUSOVÝCH DAT JAKO ZRCADLO REALITY

Velký a stále rostoucí objem jazykových dat naráží na dva krajní přístupy. Jeden, daný dnešními miliardovými korpusy (jako např. korpusy *WaCky* pro angličtinu a několik dalších jazyků, viz 3.3.5), vede k neopodstatněnému přesvědčení, že na webu je informačně dostupné vše a že se v takovém nadbytku textů najdou všechny typy jazyka a textů (nikoho pak už nezajímá otázka výběru z nich, ani jejich kvantitativního poměru). Druhý přístup, typický pro principiálně a cíleně budované reprezentativní korpusy, si však klade otázku, v jakých kvantitativních, popř. i kvalitativních (zvl. obsahových a žánrových) proporcích korpus složit. Při prvním, webovém přístupu se nikdo neptá, do jaké míry jsou texty vybírané z webu náhodné (náhodné jsou) a co v nich není (třeba spontánní mluvený dialogický jazyk), při druhém je však třeba se ptát, jak reflektovat, odrážet jazyk užívaný různými uživateli tak, aby se v korpusu zachytil v objektivně zjištěitelných a relativně stabilních proporcích. Nelze totiž předpokládat, že proporce typů jazyka jsou stále stejné, a už vůbec ne to, že webové texty, které jsou v zásadě psané, zachycují dobře i jazyk mluvený. Dnes víme, že např. proporce beletrie či poezie obecně stejné nezůstávají a nebudou ani nadále zřejmě stejné; podobné je to s rozsahem textů žurnalistických, publicistických apod. Navíc tyto a další korpusové texty by se daty získanými z webu do jisté míry paralelně zdvojovaly (pokud by se zařadily do obecného korpusu), protože se opakují (a např. zprávy o denních událostech v různých novinách se svým způsobem přímo opakují v důsledku stejných několika zdrojů, tj. tiskových agentur).

Jedním z dalších nestudovaných problémů je tedy *proměňování* žánrů a oblastí v čase. Dnes existuje oproti minulosti např. mnohem méně tištěných novin, protože jsou nahrazovány novinami elektronickými, zřejmě se také méně čte tradiční knižně vydávaná beletrie (jakkoliv ji do jisté míry vyvažují elektronické knihy). Základní předpoklad bližšího poznání je, že sociologickými výzkumy se takové proporce zjistit nejen dají, ale pro některé účely využití korpusu zjistit musejí, např. pro tvorbu velkého obecného slovníku, v nichž odpovídající žánrovou příslušnost můžou naznačovat dodané frekvence hesel. Pokud přijmeme přirozenou hypotézu, že slovník musí zachycovat (v určitém výběru a realistických proporcích) celý obecný lexikon jazyka (jakkoliv takový výběr však bude silně omezený u odborné složky jazyka), je třeba příslušné proporce typů jazyka znát. Největší část odborného jazyka, resp. jeho termínů, které nepatří do obecného jazyka a úzu, však tak jako tak zůstane mimo obecný, nespecializovaný korpus, i když je velký. Obecně však v jazyce jako celku, obecném i odborném, resp. jeho slovníku, je terminologie však jeho složkou největší.

Reprezentativnost, pokud je založená na objektivně zjišťovaných proporcích typů textů korpusu, je *jediný spolehlivý a nenáhodný základ pro výzkum* (na základě sociologického výzkumu především recepce jazyka), odrážející nutně vždy podobu jazyka v daném období. Je obecným předpokladem a principem, který je třeba

sledovat i jinde, u diachronního korpusu, mluveného korpusu, stejně jako paralelního korpusu. V každé z těchto dalších korpusových oblastí, které bude třeba postupně propojovat, se ovšem naráží na jiné faktory, které modifikují, popř. komplikují podobu reprezentativnosti korpusu. Necháme-li stranou pro diachronní korpus samozřejmou neexistenci řady žánrů, je hlavním problémem (objektivní) ne/dostupnost, resp. ne/existence určitých textů. Ta je však odlišná jak u mluveného, tak i u vícejazyčného paralelního korpusu a reprezentativnost má u nich specifická omezení.

Míra, v níž korpus věrně a proporcionálně zachycuje jazyk, se dá v korpusové lingvistice vyjádřit jako *reprezentace jazykové reality korpusem*. Na sociologickém podkladu založená se tato zjištěná míra, v podobě volených proporcí jednotlivých částí korpusu, zkoušená v minulosti především Českým národním korpusem, musí dnes vyrovnávat s novými problémy, danými mj. posunem v povaze čtenosti. Nastoluje to otázku i potřebu hledání objektivních, skutečných a současných **proporcí**

- (1) tradičního obecného psaného jazyka (tj. především publicistiky a beletrie),
- (2) jazyka mluveného a psaného (i když proporce neznáme, intuitivně víme, že mluvený jazyk je dominantní, lidé prostě stále víc mluví a poslouchají, než čtou a píšou),
- (3) obecného a specializovaného a ne všemi užívaného odborného jazyka terminologického, a konečně aspoň ještě
- (4) složek jazyka webu, v některých ohledech zvláštního (mj. svou nestabilitou a časovou anonymitou).

Jenom skrze statisticky zjištěné proporce především těchto typů textu lze podle vnějších objektivních kritérií odpovědně stavět moderní korpus a uvážený výběr z nich pak může principiálně odrážet, pokud jde o vzorek obecný, povahu celého jazyka. Jde tedy o poměr tří odlišných pojmů a entit:

(užívaný) jazyk – korpus – web.

V žádném korpusu není a nemůže být zachycený *celý jazyk*, protože ten je prakticky i teoreticky vlastně neohrazený, ale jen takto je možné, tj. zachycením proporcí jeho skutečného úzu, ho reprezentativně zachytit a aproximativně se mu přiblížit v *korpusu*, chápaném jako určitý vzorek a obraz jazyka. U webu je tato možnost nereálná, web se stále proměňuje a základní potřebu reprodukovatelnosti a citovatelnosti údajů přímo ze zdrojů většinou nesplňuje (srov. 3.3.5).

Protože se spojení jazyka psaného a mluveného nezdá zatím realistické, řešením je nechávat prozatím mluvené korpusy izolovaně stranou, protože proporce jazyka mluveného a psaného spolehlivě neznáme. V praxi je zatím bohužel pravda to, že textů psaných je k dispozici mnohem víc než mluvených, jakkoliv základním

modem komunikace je pro většinu lidí jazyk mluvený. Konečným cílem výstavby optimálního korpusu však zůstává ve správných proporcích korpus složit jak ze složky psané (která je i teoreticky menší), tak mluvené (která by měla svým rozsahem převažovat). To je však ještě daleko.

Občasná námitka tvrdící, že korpus jako (velký) vzorek není reprezentativní pro celý jazyk, protože nemůže zachytit skutečně celý jazyk ve všech jeho složkách, je málo relevantní a jen spekulativně teoretická. Lze proti ní postavit jinou námitku tvrdící, že přestože neznáme celý jazyk, ani jeho tradiční předkorpusové výzkumy nelze odmítat, protože o jazyce stále mnohé vypovídají; navíc však ani tehdy nebylo autorům zcela lhostejné, z jakých dat se vycházelo. I když řada výsledků nebyla statisticky ideální, byla často jen přibližná a orientační, ale svou cenu měla a má, což je fakt, ke kterému se leckterý dnešní statistik nechce ke své škodě znát.

Je jistě pravda, že jsou formy jazyka (dokumenty tajné, vysoce privátní, některé mluvené i dětský jazyk aj.), ke kterým není přístup vůbec (ohled ne/zveřejnitelnosti), nebo je velmi malý (ohled ne/dostupnosti), jejichž absenci je nutné při výzkumu reprezentativnosti deklarovat. Ovšem potřeba a možnost tvorby reprezentativního korpusu se takovým poukazem neanuluje, byla by to námitka mírně řečeno podivná. Nelze tu pro proklamovanou čistotu přístupu vylévat s vaničkou dítě, nesplnitelný nárok rigoróznosti některých statistických přístupů (co do možnosti pochytit z jazyka odpovědně vše, co v něm je) tu řešení situace příliš nenapomáhá. *Reprezentativní korpusy (psané) je třeba chápat jako soubory zveřejnitelných a dostupných textů, které jsou proporčně adekvátní proporcím zjištěných typů jazyka. Ty známé druhy jazyka, které do něj nemůžou vstoupit, se musejí přitom jasně označit a deklarovat a o jejich eventuální pokrytí lze usilovat až časem.*

Námitka zmiňovaná výše může nabýt i takovou podobu, že „údajná“ nereprezentativnost korpusu je dána tedy nemožností podobu a proporce korpusu vztáhnout k jakési ideální populaci textů. Má v praxi malou relevantnost už kvůli tomu, že celý jazyk skutečně nejen neznáme, ale znát nebudeme (což nás však nezabavuje potřeby ho zkoumat co nejlépe). Sociologicky založené vzorky jazyka nejsou o nic méně spolehlivé než vzorky něčeho jiného zkoumající jevy i mimo jazyk a jejich jen zprostředkovaná objektivita je stejná. Je to podobné jako při výzkumu veřejného mínění, kdy (mnohem menší) vzorky dotázaných se srovnávají s populací, kterou neznáme nijak lépe (jen podle výzkumů sociologů), a přesto se zkoumají a vychází se z nich (s naznačením procenta chybovosti). Postoje takových statistiků se tu proto jeví pošetilé a nemůžou vyvrátit praxi. Musejí však na problém kvality reprezentativnosti působit jako stimul tak, že vedou ke zlepšení. G. Leech tu výstižně mluví o „ideální“ reprezentativnosti korpusu, která (? zatím) dostupná není, její nedostatek však nemůže bránit tvorbě a výzkumu korpusů aspoň v té podobě, v jaké jsou; „zítra“ můžou být lepší.

Jeden takový uvažovaný, resp. realizovaný výzkum reprezentativnosti se opírá o míru čtenosti textů. Té se v poslední době obecně a zvl. z literárních hledisek věnuje J. Trávníček (2011, 2013) a zahrnuje už i web. Pouze čtenosti v zásadě jen periodického tisku (novin a beletrie) se věnuje Media projekt Čtenost tisku (www.rocenkaunievydavatel.cz).

Dodejme, že informace, které nabízí zvláště Google, uživatele zatím spíše svým množstvím zahlcují a záruka objektivního proporčního výběru z nich je malá. Pro specifické informace *lingvistické* v různých jazycích (ne nutně s ohledem na reprezentativnost, ale spíše na vzorky dat) je vhodnější spíše *WebCorp*, *Linguists' Search Engine* (<http://www.webcorp.org.uk/live/>).

Reprezentativní korpus *není tedy dán prostou pestrostí* zastoupení všemožných žánrů a typů textu a jejich mechanickým naskládáním vedle sebe (aniž se cokoliv tuší o jejich vzájemném poměru), *je dán pouze věrným odrazem* *proporcí skutečného úzu* (celého jazyka). Někdy se tu také ve stejném smyslu celkem zbytečně mluví o *vyváženém* korpusu (kdy vyvážení mechanicky, a tedy arbitrárně stejných proporcí jazyka proti sobě, pokud je tak míněno, o jazyce nic neříká, ničemu neodpovídá a je umělé).

Řešení reprezentativnosti korpusu, který pak nabídne objektivní a věrné zrcadlo celkového úzu jazyka na rozdíl od korpusů oportunních, a tedy náhodných, vyžaduje ovšem spolehlivé vymezení jazykových variet, žánrů, domén apod. Ty však přesně vymezené vždy nejsou a je tu mnoho překrývání a splývání, což představuje další problém, který lze však statisticky, mj. na vzorcích, řešit také. Na druhé straně to ale vyžaduje i zjištění kritérií oddělujících jazyk obecný od odborného. Má-li tedy být reprezentativní korpus věrným odrazem i obrazem jazyka, musí v maximální míře odpovídat jakožto vzorek celku jazyka. Jde-li mj. o to, dospět k věrnému popisu jazyka (zvl. v obecných slovnících a příručkách), je nutně třeba pro tento cíl mít také jeho spolehlivý reprezentativní vzorek v podobě reprezentativního korpusu. Je zatím zřejmé, že lingvisté neumějí zcela spolehlivě stanovit proporce ani jazyka synchronního a diachronního, ani jazyka psaného a mluveného (a proto se zatím studují zvlášť).

Reprezentativností se však může rozumět, vedle obecných aspektů sociolinguistických (ani ty však nelze všechny v jemnostech spolehlivě zachytit, jako je např. jazyk podle pohlaví, věku), ve zúženém pohledu i vyvážené zastoupení výrazných jazykových rysů, zvl. gramatických kategorií apod. Je to přístup, který zvlášť pro nelingvisty nemusí být příliš zajímavý. Vyžadoval by, pokud by se měl realizovat pro každý aspekt jazyka (nejen gramatiku), prozkoumat, jak a kdy je výsledný a specificky sestavený takový korpus z takového specifického a úzkého pohledu reprezentativní a co vlastně pak výzkum o jazyce vypovídá (nikdo přitom nedokázal, že gramatické kategorie a jejich proporce jsou nějak zvlášť výhodnější

než aspekty sociolingvistické). Je to jistě možné, jakkoliv pracné, přičemž případně množství takto sledovaných rysů je nedohlédnutelné.

Národní korpusy jsou korpusy jednojazyčné vztahující se (obvykle) k dominantnímu jazyku dané komunity.

2.3 KORPUSOVÝ KONTEXT

Význam slova, jak ho známe ze slovníku, které v něm významů může mít vedle sebe i víc, se v běžné komunikaci i v korpusu ujasňuje až kontextem, tj. jeho vztahy a vazbami na další relevantní slova v jeho okolí.

Prosté kombinaci *hodil mu rukavici* nemáme možnost bez kontextu jednoznačně porozumět a je třeba se podívat na širší okolí této kombinace. Ta mívá významy dva, doslovný a frazémový. Z omezeného korpusového dokladu, např. *Podle agentury AP tak **hodil rukavici** nejen Izraeli, ale i ...* se sice dovídáme, že jde dnes obecně o běžnější význam druhý („odvážnou výzvu ke střetu, ne nutně vojenskému“). Ale až dalším rozšířením kontextu zjistíme, koho se vlastně ono „házení“ týká, srov. jeho pokračování ... *americké ministryni zahraničí Condoleezy Riceové*. Že však oním činitelem je Hizballáh, se dozvíme, až když postoupíme v kontextu korpusu ještě dál, o dalších dvacet slov vlevo (SYN2010).

Terminologicky pojem *kontext* kolísá a chápe se jen volně; někdy se do něj zahrnuje i **kontext mimojazykový**, resp. **situace** (jak ho často zachycuje termín *konsituace*), ten se jindy však pojmenovává zvlášť jako kontext situace; objevuje se proto i snaha vydělit zvlášť kontext textový od mimojazykového pod alternativním názvem *co-text* (kotext).

Korpusový kontext je okolí daného slova, prvku v textu, které je při lingvistické analýze třeba k určení konkrétního významu a/nebo funkce takového slova; je podle povahy sledovaného jevu a potřeby různě dlouhý. Bez dostatečného kontextu, v příliš malém kontextu, nebo dokonce při izolovaném výskytu slovu nemusíme rozumět (viz příklad výše). Jen zřídka je už i psaný text jasně zasazený do širšího známého dobového kontextu tak, že ho každý zná a rozumí mu, např. skrze četbu novin apod., srov. slavnou úvodní větu z Haškova Švejka „*Tak nám zabili Ferdinanda*“, které rozumíme i bez jazykového kontextu knihy, protože víme, kdo byl Ferdinand aj. Většinou se ale bez širší znalosti plného a tedy jak jazykového, tak i mimokorpusového kontextu, který kontext jazykový doplňuje, avšak k plnému pochopení nestačí, neobejdeme.

V korpusu se však zpravidla dosud nijak nezachycují aspekty **kontextu mimojazykového**, mluvní situace apod., tj. všechny rysy nepsané, od gestikulace,

grimas, postojů těla až po okolní relevantní děje a zvuky. Něco málo lze také doplňkově zjistit z doprovodné korpusové databáze, znalosti autora, doby apod. Toto vše ovšem doprovázejí důležité **paralingvistické aspekty** patřící k samotnému jazyku, jako je síla, tón hlasu a intonace. U korpusu psaného, textového toto však ani dobře není možné, není to součástí psaného jazyka.

O to se do jisté míry, zatím však spíše primitivně a už za hranicemi korpusu psaného, snaží i korpus multimodální (viz 3.3).

Mají-li korpusy psané jistou, ale nikterak jednoduchou možnost doplňovat vlastní kontextovou informaci ještě odjinud (viz výše), stojí v tomto smyslu proti nim korpusy mluvené, které jsou však jak obtížnější na zpracování, tak ale také přínosnější na množství a druh informace, kterou nabízejí.

V praxi se však podstata kontextu obvykle redukuje na lineární povahu textového řetězce, tj. na kvantitativní aspekty délky okolí, a tedy (počtu) slov kolem hledaného tvaru v řádku konkordance (před i za ním), srov. už i 2.1.1. Kromě empirických postřehů žádná pevná zásada volby podoby a délky kontextu v korelaci s povahou hledaného tvaru není (srov. však Cvrček 2013). Nejvýznamnější obvykle bývá kontext bezprostřední, tedy v těsném sousedství vlevo či vpravo od sledované formy. Takto si zřejmě např. velká část adjektiv sice zhruba vystačí s kontextem následujícího substantiva a takový kontext tedy stačí obvykle i uživateli. Naopak (některé) partikule obvykle vyžadují kontext aspoň celé jedné věty, avšak pevného není dáno nic a různé typy zkoumaného tvaru vyžadují různé, resp. různé velké kontexty. Otevřenou otázkou je například, jaký má být kontext k celým větám. Některé izolované konvencionalizované věty (zvl. formule, pozdravy) jsou sice pragmaticky jasnější a vyhraněné, ale o běžných větách uvnitř větších výpovědí toho víme konkrétního jen málo. Takových vět může být i většina, a zaslouží si další zkoumání, i kromě známých implikací spojených s aktuálním členěním věty a textu. Přístup je tu tedy pragmatický, ad hoc a závisí na rozhodnutí zkoumajícího; v korpusu se ovšem podle potřeby dá snadno dohledat i kontext značně velký.

Korpusový kontext je naprosto primárním zdrojem informací a základem každé jazykové analýzy, ať formální, či sémantické, a je tedy třeba s ním pracovat citlivě, především ve smyslu způsobu jeho vyhodnocování (jeho postačitelost lze modifikovat jiným jeho zadáním, rozšířením apod.). Výhodou konkordance je možnost nashromáždění všemožných kontextů hledaného slova. Ty také napovědí, kde je centrum jeho úzu (podporované nejvyšší a vysokou frekvencí, resp. opakováním kontextu) a kde jeho periferie (ojedinělé a řídké kontexty), včetně archaismů či neologismů.

Kontext je nejenom nejvyšším *arbitrem* při hledání a posuzování konkrétní povahy na jedné straně formy, ale na druhé straně i významu a úzu

studované jazykové jednotky. Je však arbitrem i tenkrát, když taggery, taggovací algoritmy (nástroje značkování, viz 2.5.1) nejsou dostatečně přesné a narážejí mj. na homonyma. Pak je zapotřebí **disambiguace** (též *desambiguace*), rozlišení, resp. identifikace konkrétní funkce, formy nebo významu podle dostatečného kontextu (viz zvl. 3.1.3.1). V některých případech, kdy kontext nestačí, může konzultačně napomoci informace z jiných zdrojů (např. webu), ale i v konečné instanci to bývá vlastní intuice zkoumatele, která rozhoduje (a která z jeho paměti čerpá relevantní kontexty minulé).

Podle zkušenosti i teorie se obecně a v souladu s intuicí má za to, že **variabilita** kontextu se značně zmenšuje v těsné blízkosti sledovaného slova. Je to obecně dáno sémantikou slov, která se kombinují podle svého významu, resp. sémantických pravidel, a ta význam jednotlivých slov přirozeně propojují (viz mj. víc Cvrček 2013).

2.4 KORPUS A EMPIRIE

V moderní době je jazykový korpus zřetelně hlavním zdrojem studijně orientovaného poznání **obecné jazykové empirie**, která se do něj promítá jakožto úzus mnoha uživatelů a kterou jsme jen do určité míry (většinou malé) někdy schopni vyvažovat empirií a jazykovou zkušeností osobní (obvykle však jen jako korektiv u velkých výkyvů v úzu). Korpus je takto zdrojem k realizaci empirického zkoumání jazyka v protikladu k takovým přístupům, které empirii odmítají (zvl. generativismu v důsledku postojů N. Chomského). Empirismus je obecně vědecký přístup vycházející z objektivních dat, obvykle masových, a informace v nich, který stojí v protikladu k racionalismu, který se naproti tomu opírá o osobní, a tedy i nutně subjektivní introspekci. Korpusové poznání je velmi pevně a mnoha vazbami (skrze velké množství svých textů) vázané na reálná data, je takto nezpochybnitelně *vázané na realitu* (resp. v jen málo zprostředkované míře) a na roli jazyka v ní; druh poznání založený takto na korpusu je proto v maximální míře *objektivní*.

Pokud se zvláště reprezentativní korpus jako opakovaně zkoumaný zdroj etabluje pro uživatele jako respektovaný zdroj poznání, a stává se tak i cenným referenčním zdrojem, stává se z takového korpusu i model k rozšířenému studiu jazyka. Vzniká tak ale i entita nabízející se k jazykovým experimentům a dalšímu zkoumání různého druhu. *Opakovaný experiment*, základní požadavek každé vědy, se právě zde, v lingvistice kvůli „nadbytku dat“ v zásadě poprvé, nabízí v podobě korpusu jako možnost tedy také. V důsledku omezených zdrojů však opakovaný experiment narážel v minulosti na velké překážky.

Každý uživatel jazyka má za sebou svou vlastní specifickou jazykovou empirii, chápanou jako úhrn všeho, co jazykově prožil, vnímal a vyprodukoval, resp. tu