

Varitext

Varitext. <http://syrah.uni-koeln.de/varitext/> (Last Accessed: 2017-08-15). Reviewed by VIAF Julia Burkhardt (University of Leipzig), jburk [at] rz [dot] uni-leipzig [dot] de.

DOI: 10.18716/ride.a.6.3

Abstract

This review discusses the web-based platform *Varitext* that up to now provides free-of-charge access to the *Corpus des variétés nationales du français (CoVaNa-FR)* and is designed to assimilate further linguistic corpora of other languages in the future. The aim of *Varitext* is to make available large-scale resources of so-called 'pluricentric' languages focussing on the national variations of the respective standard languages. At present French standard varieties in CoVaNa-FR are represented by texts of the francophone daily press. *Varitext* offers a working environment making it possible to analyse the corpus of standard varieties with primary regard to lexical and statistical aspects.

Varitext.PrimeStat 1.0

Varitext

[Accueil](#)[Se connecter à la base](#)[S'inscrire](#)[Ressources externes](#)

Bienvenue sur le site du projet *Varitext*

Présentation

Le projet *Varitext*, animé conjointement par l'Institut des Langues Romanes de l'Université de Cologne (ILRC) et l'UMR LDI de l'Université Paris 13, s'adresse à tous les enseignants-chercheurs, doctorants et étudiants qui s'intéressent à l'analyse outillée des variétés nationales de différentes langues telles que le français. Il met à la disposition de la communauté scientifique une plate-forme d'analyse qui permet l'extraction de concordances, le calcul de spécificités fréquentielles et de cooccurrences lexico-syntaxiques.

A présent, les corpus textuels réunis dans la base *Varitext* couvrent une aire géographique importante dans l'espace francophone africain et européen.

Pour vous connecter à la base textuelle, veuillez cliquer **ici** ou sur l'onglet correspondant.

Pour vous inscrire à la base textuelle, veuillez cliquer **ici** ou sur l'onglet correspondant.

Die webbasierte Plattform *Varitext*.

Das Projekt *Varitext* hat zum Ziel, Textressourcen, und zwar konkret linguistische Korpora, sowie eine online zugängliche Arbeitsumgebung für deren Analyse bereitzustellen. Dabei sei unter Textressourcen allgemein jede beliebige Art von digitaler Textsammlung verstanden, unter einem linguistischen Korpus hingegen

speziell eine zu sprachwissenschaftlichen Zwecken zusammengestellte und aufbereitete digitale Textsammlung (s. u.). *Varitext* ist zugleich jene webbasierte Arbeitsumgebung (*plateforme*), die Zugang zu den Korpora und den damit verknüpften Analysetools verschafft (siehe Abb. 1).

Ins Leben gerufen und getragen ist das Projekt laut Internetseite vom Romanischen Seminar der Universität zu Köln und vom *Laboratoire Lexique, Dictionnaire, Informatique* (LDI) der Universität Paris 13. Hauptsächlich gestaltet und betreut wird *Varitext* von Sascha Diwersy. Zielgruppe sind laut Webseite Forschende, Lehrende und Studierende, die sich für die Erforschung und Analyse sprachlicher Varietäten interessieren. Zum Entstehungskontext und zu den Bedingungen, unter denen das Projekt realisiert wird und wurde, werden leider keine Angaben gemacht. Ebenso finden sich bedauerlicherweise keine Angaben dazu, in welchem konkreten finanziellen und institutionellen Rahmen das Projekt weiter verfolgt wird. Beschreibungen, Hilfstexte und Arbeitsoberfläche sind bislang ausschließlich in französischer Sprache gehalten. *Varitext* ist einfach und frei online erreichbar und verwendbar; es genügen eine Registrierung und die Bestätigung durch das Team hinter *Varitext*.

Das Fernziel von *Varitext* ist es, Korpora zu (Standard-)Varietäten unterschiedlicher (und durchaus nicht nur romanischer) Sprachen zu sammeln und für die Analyse zugänglich zu machen: „As is indicated by its name, it is open to host corpora for other languages compiled according to the same rationale of large scale variationist research in a pluricentric perspective.“ (Diwersy 2014, 51). Bisher bietet die Plattform Zugang zu einem französischen Korpus, dem CoVaNa-FR: *Corpus des variétés nationales du français*. Die Integration eines Korpus spanischer Texte ist Diwersy (2014, 51) zufolge in Arbeit und die Zusammenstellung weiterer Korpora (Portugiesisch, Russisch, Arabisch) zumindest geplant. Zu Anpassungen der Plattform, z. B. sprachlicher oder technischer Art, die durch eine solche Erweiterung durch andere Sprachen notwendig werden können, sind derzeit noch keine Informationen verfügbar.

Über den genauen Inhalt des Korpus CoVaNa-FR, die Aufbereitung, die Annotation und die linguistische Einordnung des Korpus informiert ein Aufsatz Sascha Diwersys (2014) (leider nicht *Varitext* selbst, s. u.). CoVaNa-FR stellt ein schriftsprachliches Korpus dar und setzt sich (bislang) aus Texten der frankophonen Presse in Frankreich, Kanada, der Schweiz sowie mehreren afrikanischen Staaten zusammen, und zwar: Elfenbeinküste, Marokko, Algerien, der Demokratischen Republik Kongo, Mali, Kamerun, Senegal und Tunesien. Geplant ist eine Erweiterung um andere schriftsprachliche, nämlich fiktionale und akademische Texte. Mit einer solchen Ressource wird in der Tat, wie Sascha Diwersy (2014, 48) meint, ein wichtiges Desiderat innerhalb der Dokumentation und Erforschung des Französischen angegangen. Zwar ist das Französische durch diverse Korpora dokumentiert:^{*} so bietet die Ressource *Frantext* z. B. einen „panchronen“ (Pusch 2014, 183) Zugang zu französischen Texten mehrerer Jahrhunderte, bei denen es sich im Kern um literarische, essayistische, philosophische und wissenschaftliche Werke handelt; zum Teil lässt sich die Textsammlung als linguistisches Korpus im engeren Sinne nutzen. Eine Vielzahl größerer und kleinerer Projekte dokumentieren unterdessen das gesprochene Französisch, wie z. B. *Corpus de Référence du Français Parlé* oder *Corpus de Langue parlée en interaction* (CLAPI). Ein umfangreiches linguistisches Korpus zur Variation des (Standard-)Französischen im internationalen frankophonen Raum, das zudem einfach und kostenlos zugänglich ist, gibt es allerdings bislang, soweit ich sehe, nicht (siehe z. B. den Überblick von Pusch 2014).

Aus fast jedem der angegebenen frankophonen Staaten sind mindestens zwei verschiedene Angebote der nationalen Tagespresse vertreten (z. B. *Le Temps* und *La Tribune de Genève* für die Schweiz, *Fraternité Matin* und *Notre Voie* für die Elfenbeinküste). Erhoben wurden für jede Zeitung jeweils alle Ausgaben eines ganzen Jahres oder mehrerer Jahre. Schwerpunktmäßig sind die Jahre 2007 und 2008 vertreten, das heißt, dass fast jedes Presseteilkorpus^{*} die Ausgaben dieser beiden oder eines dieser beiden Jahre umfasst.

Insgesamt reicht die Zeitspanne, aus der Texte erhoben wurden, von 2003 bis 2009. Presse- und Ausgabenteilkorpora können also in dieser Hinsicht als vergleichbar angenommen werden. Derzeit umfasst CoVaNa-FR ungefähr 419.700.000 Token. Die Größe der Presseteilkorpora variiert zwischen 18,8 Mio. (Elfenbeinküste) und 53,5 Mio. Token (Kanada) (Diwersy 2014, 49), in Abhängigkeit von der Anzahl der verzeichneten Jahrgänge und sicher auch vom jeweiligen Angebot der Tageszeitungen.

Laut Diwersy (2014, 51) liegt das Korpus in XML vor und ist nach Subkorpus, Einzeltext, Absatz und Satz strukturiert. Da mit der *Open Corpus Workbench* gearbeitet wird, sind die Daten vertikalisiert, also in einer Art Tabelle mit XML-Tags, formatiert. Inwieweit die Annotationen kompatibel mit den Empfehlungen der *Text Encoding Initiative* sind, wird nicht expliziert. Das Korpus ist mit morphosyntaktischen (*Part-of-Speech*) und syntaktischen Informationen (*Parsing*) angereichert und darüber hinaus lemmatisiert. Die Aufbereitung der Texte ist mit Hilfe von *Connexor Tools* (Tagger) realisiert worden. *Varitext* liefert eine benutzungsfreundliche graphische Oberfläche und umfangreiche Auswertungsoptionen, für deren Aufbau die *Open Corpus Workbench*, *UCS toolkit* 0.6 sowie *R* verwendet wurden (Diwersy 2014, 49-52). Die Art der Aufbereitung und Strukturierung der Daten erlaubt eine differenzierte, individuelle Korpuskomposition und vielfältige Analysen und Auswertungen. Die Arbeitsumgebung ermöglicht Konkordanz- (KWIC-Konkordanzen), Frequenz- und Kookkurenzanalysen (s. u.).

COVANA-FR - Corpus des variétés nationales du français

Choix du corpus

Choix des pivots

Requêtes

Aide

Définir le corpus de travail [Aide]

► Créer ou Ajouter un sous-corpus

► Créer un échantillon à partitions

Valider

Sous-corpus : Test1

Sélectionner un ou plusieurs fichiers de corpus [Aide]

Presse hexagonale

Presse camerounaise

Presse canadienne

Presse suisse

Presse ivoirienne

Presse congolaise (R.D.C.)

Presse algérienne

déconnexion

(cc) 2013 - L'équipe Varitext

Zusammenstellung des Untersuchungskorpus. Wahl der Presseteilkorpora.

Zusammenstellung des Untersuchungskorpus. Wahl der Ausgabenteilkorpora.

Um Analysen durchführen zu können, wird zunächst, unterstützt von der Arbeitsoberfläche, ein Untersuchungskorpus zusammengestellt (*créer ou ajouter un sous-corpus*, siehe Abb. 2 und 3).^{*} Es kann jederzeit verfeinert und erweitert, aber leider nicht für einen späteren Gebrauch gespeichert werden. Über die Presse- und Ausgabenteilkorpora hinaus kann das Untersuchungskorpus nach einem genaueren Datum (*Date*), nach Themen bzw. Sparten (*section*) und AutorInnen (*Auteur*) präzisiert werden.^{*} Höchst problematisch ist allerdings, dass der genaue Inhalt, die Größe und der zeitliche Kontext des Gesamtkorpus sowie seiner Subkorpora innerhalb der Arbeitsumgebung *Varitext* an keiner Stelle genau beschrieben werden. Die erste Zusammenstellung eines Untersuchungskorpus gleicht demnach einer Entdeckungsreise, weil erst im Verlaufe dieses Prozesses nach und nach erschlossen werden kann, wie CoVaNa-FR im Einzelnen aufgebaut ist.

COVANA-FR - Corpus des variétés nationales du français

Eingabe und Differenzierung der gesuchten Wörter oder Ausdrücke.

COVANA-FR - Corpus des variétés nationales du français

Choix du corpus
Choix des pivots
Requêtes
Aide

Concordances KWIC
Fréquences
Cooccurrences

Définir les paramètres de tri de la concordance [\[Aide\]](#)

Trier par:
Empan:

Mot-forme
3G...1G
2G...1G
1G
1G..Pivot..1D
Pivot
1D
1D...2D
1D...3D

Afficher la concordance KWIC

Präzisierung der Anfrage. Auswahl der Analysefunktion (KWIC-Konkordanz, Frequenz, Kookkurrenz), ggf. Regulierung der Ergebnisanzeige und des Kontextes.

COVANA-FR - Corpus des variétés nationales du français - Résultats lexico-statistiques

Lexicogramme
Tableau croisé
Graphiques
Aide

Lexicogramme [\[Aide\]](#)

Afficher 10 lignes Filtre global :

Pivot	Echantillon	Fréquence	Fréquence - Pivot	Taille - Echantillon	Totaux	Occurrences	Rang Occurrences
la__DET++femme__N	Test1	3238	3238	43975946	43975946	3238	1
Pivot	Echantillon	Fréquence	Fréquence - Pi	Taille - Echantill	Totaux	Occurrences	Rang Occurrence

Ligne(s) 1 à 1 sur 1

Einfache Frequenzanalyse: Lexicogramme von *la+femme*.

Lexicogramme

Tableau croisé

Graphiques

Aide

Lexicogramme

[Aide]

Afficher

10

lignes

Filtre global :

Pivot	Collocatif	Fréquence	Fréquence - Pivot	Fréquence - Collocatif	Totaux	Echantillon	Log-likelihood	Rang Log-likelihood
la++femme::Test1	de__PREP__L	2855	16190	18086630	219879730	Test1	1469,7759	1
la++femme::Test1	marocain__A__R	226	16190	137545	219879730	Test1	975,0699	2
la++femme::Test1	promotion__N__L	153	16190	40405	219879730	Test1	907,5952	3
la++femme::Test1	droit__N__L	194	16190	102355	219879730	Test1	889,8320	4
la++femme::Test1	journée__N__L	142	16190	69465	219879730	Test1	671,5888	5
la++femme::Test1	quinzaine__N__L	76	16190	5075	219879730	Test1	658,1281	6
la++femme::Test1	de__N__R	68	16190	12115	219879730	Test1	455,8252	7
la++femme::Test1	chez__PREP__L	98	16190	62195	219879730	Test1	414,2621	8
la++femme::Test1	rôle__N__L	92	16190	53320	219879730	Test1	404,8439	9
la++femme::Test1	entreprenariat__N__R	41	16190	1865	219879730	Test1	386,5977	10
Pivot	Collocatif	Fréquence	Fréquence -	Fréquence -	Totaux	Echantillon	Log-likelihood	Rang Log-likelihood

Ligne(s) 1 à 10 sur 211

Télécharger : [.csv]

Kookkurrenzanalyse zu „la femme“, zeigt häufigste linke (_L) und rechte (_R) Kollokatoren aus und gibt entsprechend den Einstellungen in der Anfrage statistische Werte aus.

Steht das (vorläufige) Untersuchungskorpus fest, können detaillierte Suchanfragen formuliert und entschieden werden, welche Art der Auswertung (KWIC-Konkordanzen, Frequenz-, Kookkurrenzanalysen) durchgeführt werden soll (siehe Abb. 4 und 5). Die Bearbeitung der zugrundeliegenden Textbasis (s. o.) erlaubt die Suche von exakten Wortformen, Lemmata, Kollokationen und die Verwendung von regulären Ausdrücken. Suchwörter und –phrasen können morphosyntaktisch spezifiziert werden, und zwar nach Wortart und nach syntaktischen Relationen und Funktionen (z. B. *sujet*, *attribut de l'objet*, *auxiliaire*). Eine Kombination verschiedener Kriterien ist möglich und ermöglicht komplexe Abfragen. Kookkurrenz- und Frequenzanalyse bieten umfangreiche statistische Auswertungen, die in verschiedenen Formaten (Lexikogramm, Tabelle, Graphik) dargestellt werden (siehe Abb. 6 und 7).

Concordance KWIC

Aide

Table de la concordance [\[Aide\]](#)

Afficher 10 lignes

Filtre global :

Occurrence No.	Corpus	Cotexte gauche	Empan gauche	Pivot	Empan droit	Cotexte droit	Méta
1	PRESSE_FRA	: d'être exclue. Mineur et jeunesse : telles sont les valeurs absolues de la féminité aujourd'hui. Autrement dit, il est interdit de dépasser 55 kg et d'avoir des rides (et encore pire, de dépasser 55 kg en ayant des rides	!)	La femme	à 30	ans se retrouve confrontée à la difficulté de devoir mener deux vies juxtaposées. Ayant toujours des salaires plus bas que ceux des hommes à compétence égale, elle reste défavorisée. Prise en tenaille entre sa vie professionnelle et sa vie de femme, de	⊖
Sous-échantillon: LFI08 Titre: Petit bilan de la condition féminine en 2008 Date: 08/03/2008 Auteur: NULL Section: DÉBATS							
1	PRESSE_MAR	: sur " La femme artiste au Maroc et dans le monde arabe " dont la publication coïncidera avec le 8 mars prochain, a mis l'accent sur l'apport de la femme au cinéma marocain, notant que son livre retrace la progression " éblouissante	" de	la femme	marocaine dans	le domaine artistique, toutes disciplines confondues. Japon en récession Les marchés déçus par le G20 Le Japon, deuxième économie mondiale, a annoncé lundi son entrée en récession tandis que les marchés asiatiques évoluaient en ordre dispersé, beaucoup d'investisseurs étant déçus par l'absence	⊕
1	PRESSE_SEN	: Mais la Jérusalem céleste est libre et c'est elle notre mère ". En effet, l'Ecriture déclare : 'Réjouis-toi, femme qui n'avais pas d'enfants (Ndlr : Sarah, mère d'Isaac, l'ancêtre des juifs)	! Car	la femme	abandonnée (	Ndlr : Agar) aura plus d'enfants que la femme aimée par son mari (Ndlr : Sarah) (Epître de Paul aux Galates 4 : 22 - 28). C'est là une preuve de plus, si besoin est, que	⊕
2	PRESSE_FRA	: revanche, la mortalité a tendance à baisser. Cette hausse est liée à l'essor démographique et au vieillissement de la population certes, mais pas seulement. Un peu plus de la moitié des cas supplémentaires, 52 % chez l'homme et 55	% chez	la femme	, échappe	à ces paramètres. Avec au final 180 000 nouveaux cas chez les hommes en 2005 (prostate, poumon et colon-rectum pour les trois plus fréquents) et 140 000 chez les femmes (sein, colon-rectum et poumon). A u niveau des décès,	⊖

Anzeige der Ergebnisse. Ergebnisse für KWIC-Konkordanzen zu *la+femme*.

Die Ausgabe der KWIC-Konkordanzen (möglich bis maximal 10.000 Ergebnisse) erfolgt in einer Tabelle, die das Suchwort, die unmittelbaren linken und rechten Nachbarn sowie den linksseitigen und rechtsseitigen Kotext (bis zu 50 Wörter) übersichtlich präsentiert. Zudem können ggf. die jeweiligen Presseteilkorpora, aus denen die Belege stammen, nachvollzogen und weitere Metadaten abgerufen werden (siehe Abb. 8). Die Inhalte der Tabellenspalten können zusätzlich sortiert und gefiltert werden. Die Ergebnisse sind in Form einer csv-Datei herunterladbar, stehen also bei Bedarf für eine weitere Verarbeitung zur Verfügung. Ein Zugriff auf die Volltexte ist jedoch nicht möglich. Die Zugriffsbeschränkungen verschiedener Art werden nur ungenau erläutert und begründet; so finden sich im Hilfetext lediglich recht allgemeine Hinweise dazu, dass die Ausgabe von Konkordanzen oder der Umfang des Kotextes „pour des raison d'ordre juridique (et technique)“ begrenzt ist. Deutlich gesagt wird, dass aus Gründen des Copyrights CoVaNa-FR nicht als komplette Ressource heruntergeladen werden kann, sondern ausschließlich in Teilen und über die graphische Benutzungsoberfläche zur Verfügung steht (Diwersy 2014, 49).

Es können hier nicht alle Funktionen ausführlich beschrieben werden: *Varitext* ist in jeder Hinsicht ein wertvolles Arbeitsinstrument für die Analyse der französischen Sprache. Die drei genannten Auswertungsfunktionen qualifizieren es insbesondere für Fragestellungen im Zusammenhang mit dem Lexikon, die Textzusammenstellung für den Vergleich innerhalb der Frankophonie. Das Hilfemenü (*Aide*) gibt zu allen Analysefunktionen eine verständliche und detaillierte Anleitung (in französischer Sprache), und die übersichtliche, teils anschauliche Darstellung der Ergebnisse erleichtert die Untersuchung und / oder Interpretation der Daten. So können die Möglichkeiten von *Varitext* auch von z. B. Studierenden mit vergleichsweise wenig Aufwand ausgeschöpft werden. Gleichzeitig bietet die Arbeitsplattform fortgeschrittenen Nutzenden ausreichend Differenzierungsmöglichkeiten bei der statistischen Auswertung.

Problematisch, da mehr als lückenhaft, ist leider insgesamt die Dokumentation und Beschreibung der Plattform *Varitext* und des Korpus CoVaNa-FR sowie seiner Voraussetzungen, konzeptionellen Verortung und genauen Zielstellung. Dies gilt in korpusmethodischer ebenso wie in variationslinguistischer Hinsicht. Zwar hat Sascha Diwersy in seinem 2014 veröffentlichten Beitrag zu den *Proceedings of the First Workshop*

on Applying NLP Tools to Similar Languages, Varieties and Dialects Vieles zu einer solchen Dokumentation beigetragen; jedoch ist fast nichts davon auf der Internetseite selbst nachvollziehbar und ist die Publikation dort auch nicht abgelegt oder zumindest verlinkt; es findet sich auch kein Hinweis darauf.

Aus (varietäten-)linguistischer Sicht fehlt eine Beschreibung und Einordnung von CoVaNa-FR im Rahmen der Plattform *Varitext* fast vollständig. Die Kurzvorstellung (*Accueil*, siehe Abb. 1) und die Nutzungsbedingungen (*Charte Varitext*) benennen als Inhalt der Textbasis geographische Varietäten der afrikanischen und europäischen Frankophonie: „la base *Varitext* couvrent une aire géographique importante dans l'espace francophone africain et européen“*. In diesem Zusammenhang findet der Begriff ‚variétés nationales‘ Verwendung, der auf der Webseite nicht weiter expliziert wird, sich aber auch im Namen der konkreten Textbasis wiederfindet: *COVANA-FR – Corpus des variétés nationales du français*. Es ist davon auszugehen, und das bestätigt Diwersy (2014, 48-49), dass damit auf das Konzept der ‚nationalen Varietäten‘ sogenannter plurizentrischer Sprachen abgehoben wird. In dessen Sinne versteht z. B. Bernhard Pöll unter nationalen Varietäten „Ausprägungen der jeweiligen Standardvarietät in einem Gebiet, das entweder mit einem Staat oder mit einem staatsähnlichen Gebilde zusammenfällt“ (Pöll 2000, 51). Dies vorausgesetzt, bezieht *CoVaNa-FR* sich also auf diejenigen Varietäten des Französischen, die in den jeweiligen frankophonen Staaten als „internes Standardfranzösisch“ anerkannt und als solches zum hexagonalen Standard in Beziehung gesetzt werden können (Diwersy 2014, 48). Dies birgt jedoch verschiedene Probleme und sollte deshalb keinesfalls unkommentiert bleiben. Die fraglichen Punkte sollen hier angedeutet werden:

Erstens ist die Identifikation von Variations- und Nationengrenze, den der Begriff *variétés nationales* impliziert, nicht unproblematisch (siehe z. B. die Diskussion des Begriffs bei Reiffenstein 2001). Diwersy zufolge fokussiert die Korpuskonstruktion allerdings „elements of endonormative differentiation, i.e. the emergence of regionally specific norms compared to a supposed metropolitan standard variety of French“ (Diwersy 2014, 48). Ob solche „regionalspezifischen Normen“ mit dem Begriff ‚variété nationale‘ passend erfasst sind, wäre zu diskutieren. Diwersy findet hier eine Formulierung, die eigentlich für eine viel offenere begriffliche Konzeption steht und in der Lage wäre, Heterogenität innerhalb frankophoner Regionen zu integrieren. Die höchst komplexen und je eigenen Verhältnisse gerade in den multiethnischen und multilingualen afrikanischen Staaten lassen überhaupt die Einordnung des Französischen als ‚nationale‘ Sprache / Varietät in diesen ehemaligen Kolonien diskutabel erscheinen.

Zumal sich zweitens diese „nationalen Standardvarietäten“ in ihren Funktionen für die Sprachgemeinschaften und in ihrer Reichweite durchaus stark unterscheiden können (cf. Bickel 2000); Standardsprachen bilden also nicht notwendigerweise ein einheitliches und auch kein selbsterklärendes Phänomen. In den frankophonen Staaten Afrikas hängt, wie Pöll zu entnehmen ist (1998, 95 ff.), die Reichweite des Französischen, zumal des Standardfranzösischen, nicht selten von Status, Bildung und ethnischer Zugehörigkeit der einzelnen Sprechenden ab. Bisweilen ist Französisch reine Schrift- und Verwaltungssprache. Das trifft so jedoch natürlich nicht für alle frankophonen Regionen zu. Die hiesige Konzeption von Standardsprache und -norm bezieht sich aber ganz offensichtlich allein auf die geschriebene Sprache und identifiziert diese mit Standardvarietäten: „It should be obvious, then, that our present activities focus on diversifying the corpus resources, especially with regard to other written genres.“ (Diwersy 2014, 55). Dies wird in *Varitext* weder erwähnt noch begründet.

Drittens ergibt sich aus der (rudimentären) Beschreibung des Korpus auf der Internetseite einerseits und dem manifesten Inhalt des Korpus andererseits zunächst der Eindruck, dass hier stillschweigend vorausgesetzt wird, diese ‚nationalen (Standard-)Varietäten‘ würden durch die jeweilige überregionale Tagespresse repräsentiert. Das ist sicher keine abwegige Annahme, sie sollte aber dennoch begründet sein.

In seinem Konferenzbeitrag verweist Diwersy (2014, 49) auf entsprechende Überlegungen zum Verhältnis zwischen Presse und Norm und betont auch, dass ein Ausbau des Korpus um weitere Textsorten vorgesehen ist:

Due to its focus on endonormative differentiation, the CoVaNa-FR is less balanced with respect to genre than similar corpora for other languages [...] The initial version of the CoVaNa-FR, accessible on the Varitext platform, is made up of journalistic texts published by national newspapers in different Francophone countries in Africa, Europe and North America. The choice of national newspapers as primary sources is based on the assumption made by Glessgen (2007: 97) that these are particularly representative of contemporary standard varieties (...).

Diwersy 2014, 49.

Es lässt sich diskutieren, ob Presstexte Standardvarietäten „repräsentieren“ oder ob nicht vielmehr eine enge, wechselseitige Beziehung zwischen der Sprache überregionaler Presseerzeugnisse und sprachlichen Normen angenommen werden muss (siehe hierzu z. B. Burr 2004; Burkhardt et al. in Vorbereitung); und es muss gefragt werden, ob Presstexte *allein* in der Lage sind für Standardvarietäten zu stehen, auch wenn das nur vorübergehend der Fall ist. Die Bezeichnung des bisherigen Korpus oder / und dessen Beschreibung sollten diesen Einschränkungen gegebenenfalls Rechnung tragen.*

Insgesamt wäre für die Nutzung der Ressource als linguistisches Korpus im engeren Sinne die Dokumentation

- a)
- b)
- c)

unmittelbar im Rahmen der Präsentation des Korpus mehr als wünschenswert. Dies gilt ebenso für eine Beschreibung

- d)
- e)
- f)

Tatsächlich findet sich nur in den Nutzungsbedingungen von *Varitext* ein knapper Hinweis auf das XML-Format sowie auf PoS-Tagging, Lemmatisierung und Parsing. Dass das Korpus überhaupt und ausschließlich Presstexte enthält, „entdeckt“ die Nutzerin des Korpus erst, wenn sie beginnt, ein Untersuchungskorpus zusammenzustellen (s. o.).

Dass CoVaNa-FR nicht nur als linguistisches Korpus gedacht ist (dies vermittelt zumindest der Name), sondern als solches nutzbar ist, steht dabei außer Frage. Unter einem linguistischen Korpus im engeren Sinne verstehe ich in Anlehnung an Biber, Conrad und Reppen (1998, 246) sowie Burr (2004, 136-137) eine Sammlung elektronisch vorliegender Texte, die a) authentischen sozialen Kontexten entstammen, b) genau so, wie sie dort vorkamen, in das Korpus aufgenommen worden sind und c) mit Annotationen versehen sind, die mindestens die Zuordnung der Daten zu ihrem (zeitlichen, räumlichen usw.) Kontext erlauben. Ein linguistisches Korpus unterscheidet sich zudem von anderen ‚Textkollektionen‘ durch seine überlegte Zusammenstellung und Bearbeitung für linguistische Untersuchungen. In diesem Sinne definiert John Sinclair:

A corpus is a collection of pieces of language text in electronic form, selected according to external criteria to represent, as far as possible, a language or language variety as a source of data for linguistic research.

Sinclair 2005, o. S.

Das drückt sich im Falle der Textbasis von *Varitext* nicht nur in der Zusammenstellung nach externen Kriterien aus, wie es Sinclair fordert,^{*} sondern eben auch in der umfangreichen linguistischen Annotation, die konkrete sprachwissenschaftliche Analysen stützt.

Ein wichtiges Kriterium ist darüber hinaus aber auch, dass das Korpus ausgewählte und beschreibbare Sprachausschnitte enthält, die geeignet sind, eine Sprache oder einen bestimmten Teil der Sprache (z. B. eine Varietät) in gewisser Hinsicht zu repräsentieren; die sich daraus ergebende Zusammensetzung bestimmt erheblich mit, welche Fragestellungen mit Hilfe des Korpus bearbeitet werden können: „The representativeness of the corpus, in turn, determines the kinds of research questions that can be addressed and the generalizability of the results of the research.“ (Biber, Conrad, and Reppen 1998, 246). Jedoch erfordert sowohl die Beurteilung der ‚Repräsentativität‘ eines Korpus in Bezug auf eine bestimmte Varietät als auch die Einschätzung, inwieweit die Ergebnisse von Datenanalysen verallgemeinerbar sein werden, als auch die Antizipation möglicher Fragestellungen, die mit Hilfe eines Korpus beantwortet werden könnten, gerade eine transparente, umfangreiche und methodisch wie konzeptionell fundierte Dokumentation.

Die Arbeitsplattform *Varitext* und das bisher enthaltene linguistische Korpus *Corpus des variétés nationales du français* stellen eine vielversprechende Ressource für die linguistische, und zwar insbesondere lexikalische und lexikologische Analyse französischer Varietäten in Europa, Afrika und Nordamerika dar. Zum jetzigen Zeitpunkt lässt sich das Korpus dabei in erster Linie als ‚Korpus frankophoner Pressesprachen‘ nutzen. Die Arbeitsumgebung *Varitext* bietet einen vergleichsweise benutzungsfreundlichen Zugang zum Textkorpus und stellt erhellende automatische Analysen sowie statistische Auswertungen bereit. Als dringend erforderlich erweist sich eine transparentere Darstellung und Beschreibung des Projekts, der Plattform und vor allem des Korpus; dies gilt sowohl in Hinsicht auf die methodische und inhaltliche Korpuskonstruktion als auch auf die sprachwissenschaftliche Verortung seines Inhalts. Eine solche Dokumentation stellt (einmal ganz abgesehen von ihrem Interessantheitswert) aus korpuslinguistischer Sicht die Voraussetzung für eine adäquate Nutzung der Ressource und die anschließende Bewertung der Ergebnisse dar.

References

- Analyse et traitement informatique de la langue française. 2017. “Base textuelle FRANTEXT.” Accessed 30.05.2017. <http://www.frantext.fr/>.
- Biber, Douglas, Susan Conrad, and Randi Reppen. 1998. *Corpus Linguistics: Investigating language structure and use*. Cambridge: Cambridge University Press.
- Bickel, Hans. 2000. “Deutsch in der Schweiz als nationale Varietät des Deutschen.” *Sprachreport* (4): 21-27.
- Burkhardt, Julia, Elena Potapenko, Rebekka Sierig, and Aramis Concepción Durán. In Vorbereitung. “Leipziger Frltz-Korpus: Das Korpus französischer und italienischer Zeitungssprache und seine Auszeichnung.” *Romanische Studien*.

Burr, Elisabeth. 2004. "Das Korpus romanischer Zeitungssprachen in Forschung und Lehre." In *Romanistik und neue Medien: Romanistisches Kolloquium XVI*, edited by Wolfgang Dahmen, Günter Holtus, Johannes Kramer, Michael Metzeltin, Wolfgang Schweickard, and Otto Winkelman, 133-162. Tübingen: Narr.

Laboratoire ICAR. 2014. "Corpus de Langue parlée en interaction." https://web.archive.org/save/_embed/http://clapi.ish-lyon.cnrs.fr/.

Diwersy, Sascha. 2014. "The Varitext platform and the Corpus des variétés nationales du français (CoVaNa-FR) as resources for the study of French from a pluricentric perspective." *Proceedings of the First Workshop on Applying NLP Tools to Similar Languages, Varieties and Dialects*: 48-57. doi: 10.3115/v1/W14-5306.

Pöll, Bernhard. 1998. *Französisch außerhalb Frankreichs: Geschichte, Status und Profil regionaler und nationaler Varietäten*. Tübingen: Niemeyer.

Pöll, Bernhard. 2000. "Plurizentrische Sprachen im Fremdsprachenunterricht (am Beispiel des Französischen)" In *Normen im Fremdsprachenunterricht*, edited by Wolfgang Börner, Wolfgang and Klaus Vogel, 51-63. Tübingen: Narr.

Pusch, Claus D. 2014. "Les corpus romans contemporains." In *Manuels des langues romanes*, edited by Andre Klump, Johannes Kramer and Aline Willems, 173-196. Berlin / Boston: de Gruyter.

Reiffenstein, Ingo. 2010. "Das Problem der nationalen Varietäten. Rezensionssatz zu Ulrich Ammon: Die deutsche Sprache in Deutschland, Österreich und der Schweiz. Das Problem der nationalen Varietäten, Berlin/New York 1995." *Zeitschrift für Deutsche Philologie* 1: 78-89. Accessed 10.04.2017. <https://www.zfdphdigital.de/ZFDPH.01.2001.078>.

Siepmann, Dirk, Christoph Bürgel, and Sascha Diwersy. 2016. "Le Corpus de référence du français contemporain (CRFC), un corpus massif du français largement diversifié par genres." In *SHS Web of Conferences* (27): 11002. doi: 10.1051/shsconf/20162711002.

Sinclair, John. 2005. "Corpus and Text: Basic Principles." In *Developing Linguistic Corpora: a Guide to Good Practice*, edited by Martin Wynne. Oxford: Oxbow Books <https://web.archive.org/web/20170530110622/http://ota.ox.ac.uk/documents/creating/dlc/>.

Figures

Die webbasierte Plattform *Varitext*.

Zusammenstellung des Untersuchungskorpus. Wahl der Presseteilkorpora.

Zusammenstellung des Untersuchungskorpus. Wahl der Ausgabenteilkorpora.

Eingabe und Differenzierung der gesuchten Wörter oder Ausdrücke.

Präzisierung der Anfrage. Auswahl der Analysefunktion (KWIC-Konkordanz, Frequenz, Kookkurrenz), ggf. Regulierung der Ergebnisanzeige und des Kontextes.

Einfache Frequenzanalyse: Lexicogramme von *la+femme*.

Kookkurrenzanalyse zu „la femme“, zeigt häufigste linke (_L) und rechte (_R) Kollokatoren aus und gibt entsprechend den Einstellungen in der Anfrage statistische Werte aus.

Anzeige der Ergebnisse. Ergebnisse für KWIC-Konkordanzen zu *la+femme*.

Notes

Zum Französischen steht bislang kein ausgewogenes Referenzkorpus in dem Sinne zur Verfügung, wie es z.B. für das Deutsche mit dem DeReKo (Deutsches Referenzkorpus) des Instituts für deutsche Sprache der Fall ist. An einem *Corpus de référence du français contemporain* arbeiten jedoch Dirk Siepmann, Christoph Bürgel und Sascha Diwersy (2016).

Das Gesamtkorpus CoVaNa-FR besteht aus Subkorpora auf der Ebene der Länder und ihrer jeweiligen Presse (*presse hexagonale, presse suisse, presse marocaine...*). Diese wiederum bestehen aus den Subkorpora, die sich durch die Sammlung der Ausgaben eines bestimmten Jahres ergeben (z.B. *Le Monde* 2007). Ich bezeichne die Subkorpora der ersten Ebene als Presseteilkorpus. Die Subkorpora der zweiten Ebene als Ausgabenteilkorpus.

Alternativ kann für die Erstellung des Untersuchungskorpus die Variante *créer un échantillon à partition* gewählt werden. Die Varianten unterscheiden sich im Wesentlichen darin, dass die erste (*créer un copus*) einem top-down-Prozess entspricht: das Gesamtkorpus wird durch immer detailliertere Kriterien gegliedert. Unterdessen bietet die zweite einen schnelleren Zugang zu einem beliebig großen Teilkorpus, das nach einem präzisen Kriterium wie z. B. Datum gefiltert wurde.

Das ebenfalls mögliche Kriterium *Titre du Texte* liefert keine Ergebnisse. Weitere Filterkriterien sind *code géographique* und *code sous-échantillon*, da jedem Presseteilkorpus ein geographischer Code (z. B. CAN für Kanada) sowie ein Korpus-Code (z. B. PRESSE_CAN) zugeordnet ist. Jedes Ausgabenteilkorpus hat ebenfalls einen Code (z. B. DEV_CAN07 für *Le Devoir* 2007, Kanada).

Hier wurde Kanada vergessen.

Da sich das Korpus ausdrücklich auch an Studierende richtet, sollte das schon aus didaktischen Gründen geschehen.

Mit externen Kriterien ist die kommunikative Funktion von Texten gemeint, mit internen konkrete sprachliche Merkmale, die Sinclair (2005) zufolge jedoch nicht Auswahlkriterium, sondern vielmehr Gegenstand der Untersuchung sein sollen.

Licence: <http://creativecommons.org/licenses/by/4.0/>

Published: 2017-09 · build cb141fb (2026-08-23T11:45:27+02:00)