

Historical Corpus of Dutch application manual

Table of contents

Introduction	2
Information about the corpus published in the application	3
Metadata categories	3
Main	3
Period	3
Region	3
Genre	3
Details	3
Author	3
Place	3
Year	3
Permissive / Strict	4
Application user manual	5
Getting started	5
Searching the corpus	6
Simple search	6
Search	6
Wildcards	6
Reset	7
History	7
Global settings	7
Extended search	8
Starting a new search	9
Filter search by	9
Filter by year	10
Advanced search	11
The query builder	11
The tab search	11
Token attributes	12
Adding attributes to a token box	12
Function of the two +-buttons in a token box	13
The tab Context	13
Managing sequences of token boxes	14
Uploading value lists in the query builder	14
Copy to CQL editor	14
Expert search	14

Copy to query builder	15
Import query	15
Gap filling	15
Viewing results	16
Per Hit view	17
Sorting results	17
Grouping results	17
Per Document view	18
Sorting results	18
Grouping results	19
Exporting results	19
Information about a document	19
Content	19
Metadata	20
Statistics	20
Exploring the corpus	20
Documents	20
N-grams	21
Options	21
Example	22
Statistics (frequency lists)	23
Options	23
Example	23
Appendix: Corpus Query Language	24
CQL support	24
Supported features	24
Differences from CWB	25
(Currently) unsupported features	25
Using Corpus Query Language	26
Matching tokens	26
Sequences	27
Regular expression operators on tokens	27
Punctuation	27
Case- and diacritics-sensitivity	28
Matching XML elements	28
Labeling tokens, capturing groups	29
Global constraints	29

Introduction

This manual describes the corpus exploitation environment of the Historical Corpus of Dutch (HCD). The corpus is a project of the Vrije Universiteit Brussel and the Universiteit Leiden in collaboration with the Instituut voor de Nederlandse Taal (Dutch Language Institute).

The corpus application is developed by the INT. The backend of the application is the BlackLab Lucene based search engine developed for corpora with token-based annotation (<https://blacklab.ivdnt.org/>). The web-based frontend is a further development of the corpus-frontend application developed by INT (<https://github.com/instituutnederlandsetaal/blacklab-frontend>). Its design is inspired by the first version of the OpenSoNaR user interface by Tilburg University and Radboud University (<https://github.com/Taalmonsters/WhiteLab2.0>).

Information about the corpus published in the application

The Historical Corpus of Dutch (HCD) is a diachronic, regionally balanced, multigenre corpus of written Dutch. It aims to fill an important gap in the research infrastructure for historical Dutch, which has long lacked a balanced corpus with data from across the centuries and from various regions and genres.

Metadata categories

The HCD has been enriched with a set of metadata categories. These metadata will all be described below. In the corpus application it is possible to limit a search by filtering on metadata categories.

Main

Period

There are four time periods: 1530-1570, 1630-1670, 1730-1770, and 1830-1870.

Region

The HCD comprises textual material from four regions in the northern and southern Low Countries: Holland and Zeeland in the north (in the present-day Netherlands), and Brabant and Flanders in the South (in present-day Belgium).

Genre

The HCD comprises administrative texts, ego-documents, and pamphlets.

Details

Author

It is possible to search by author name.

Place

It is possible to search for Place(s) (in modern spelling).

Year

The year in which the document was written or the period in which the documents were written.

Permissive / Strict

It is possible to do a permissive or a strict search for Witness Year. If you do a *permissive* search for the years 1530-1540 you will find 8 documents, whereas you only get 6 documents if you do a *strict* search for those same years.

Application user manual

Getting started

Here are a few examples of what you can do with the corpus application (the links will take you to the application):

- To search for a heavily normalised word, use Simple Search:
 - Simple Search for Word [moeder](#)
 - Extended Search for Word [vader](#)
- To search for words satisfying a certain pattern, use *wildcards* in Simple Search or Extended Search, or *regular expressions* in Expert Search
 - words starting with *ver* and ending with *len* in [Simple Search](#)
 - words starting with *ver* and ending with *len* in [Extended Search](#)
 - words starting with *ver* and ending in *eren* with one syllable in between in [Expert Search](#)
- To see which unique forms occur as a result of your search, use Group Results.
 - example Group by Annotation: [different words following lieve](#)
 - example Group by Annotation: [different words preceding the word huis](#)
- To explore the distribution of document properties in the corpus, use the Explore feature
 - example: [characteristics about the regions of the documents](#)
 - example: [characteristics of the author of the documents](#)

Searching the corpus

Simple search

Search

The Simple Search allows you to quickly search for specific word forms (e.g. *huis*). After entering a search term, a spinner briefly appears on the right side of the search bar. Based on the keyed in word, suggestions are given of possible variants of spelling and/or form from the [GiGaNT-lexicon](#).

Based on the information in this lexicon all spelling variants of the search term found are suggested (see the screenshot below). You can then choose from the presented suggestions or select all at the same time (Select all). To make your search even more targeted, it is also possible to limit the search to the parts of speech that were found in the historic component of the GiGaNT-lexicon in connection to the search term.

If you know exactly which word you are looking for, you can also – while the wheel is spinning – press Enter directly. The search will then start immediately.

It is also possible to enter a phrase: *vrauwen ende kinderen* or *zoo als ik*. You will then find all occurrences of that exact phrase.

Note that in Simple Search the patterns will be matched case-insensitively: *god* for instance will deliver the same results as *God* or *GOD*. See the paragraph [Expert Search](#) to see how it is nevertheless possible to distinguish between uppercase and lowercase letters.

Wildcards

In Simple Search, the use of wildcards can prove good service to search for specific word forms or lemmata. A wildcard is a symbol used to replace or represent one or more characters. The following two wildcards are supported:

- * The asterisk matches any character zero or more times. Therefore, searching for *a*n* matches all word forms that start with an *a* and end with a *n*, e.g. *aen*, *allen*, *alleen*, *Antwerpen* and so on.
- ? The question mark matches a single character once. Therefore, searching for *b?n* matches *only* three-letter word forms or lemmata starting with an *b* and ending with a *n*, e.g. *ban*, *ben*, *bin*, *bon*.

This wildcard can be used more than once. Thus *d???n* matches *dijen*, *desen*, *dagen* and *deden* and so on.

Note that searching with wildcards is limited to Simple Search and Extended Search. [In Advanced Search and Expert Search you can use so-called regular expressions instead of wildcards.]

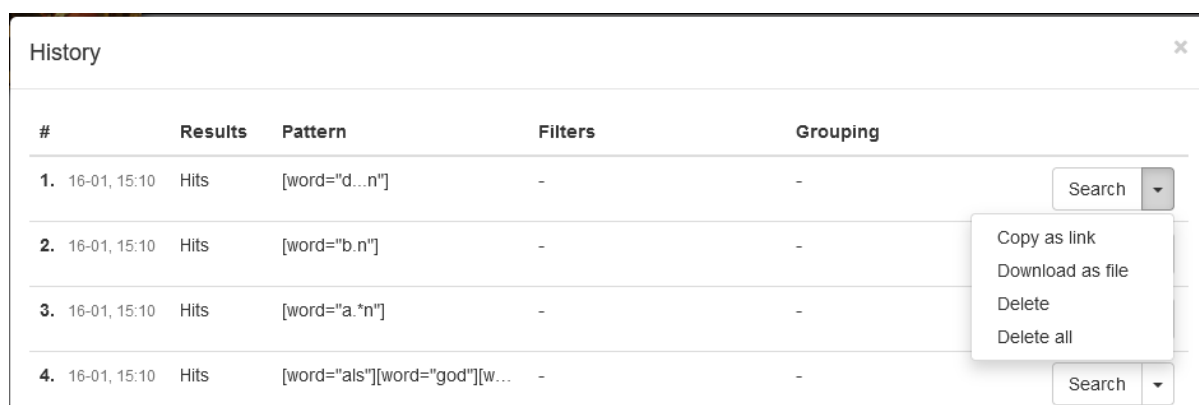
Reset

You can start a new search by pressing the Reset button. By doing so, both the search query and the hits found will be cleared. Your search history, however, will remain unchanged.

Note that it is also possible to start a new search by entering a new word or phrase in the search field.

History

The History button will display your query history. Per search query there are several possibilities (as shown in the screenshot below): you can perform the search query again (Search), you can copy the search query as a link (Copy as link), you can download the search query as a file (Download as file), you can delete a single search query (Delete) or delete all search queries (Delete all).



The screenshot shows a 'History' panel with a table of search queries. A context menu is open over the table, showing options: Search, Copy as link, Download as file, Delete, and Delete all. The table has columns: #, Results, Pattern, Filters, and Grouping. There are four rows of search history.

#	Results	Pattern	Filters	Grouping
1. 16-01, 15:10	Hits	[word="d...n"]	-	-
2. 16-01, 15:10	Hits	[word="b.n"]	-	-
3. 16-01, 15:10	Hits	[word="a.*n"]	-	-
4. 16-01, 15:10	Hits	[word="als"][word="god"][w...	-	-

Every search query has its own url. If you copy this url via History (Copy as link) or directly from the address bar of your browser, you can send it to someone else who can import this link via Import from a link. It offers that person the possibility to run the search on his own computer.

Global settings

The Global settings dialogue, activated by pressing the wheel button, allows you to configure five settings: Results per page, Sample size, Seed, Context size and Wide View.

- *Results per page*: you can choose whether you want 20, 50, 100 or 200 results to be shown;
- *Sample size*: selecting a value here will instruct the search engine to return a random sample drawn from the complete result set. The sample size can be limited by
 - a percentage of the total number of search results (percentage);
 - the number of results displayed (count).
- *Seed*: a 'random seed' is a number used to initialise a so-called pseudo-random number generator. Keeping the same seed will ensure that two samples drawn from the same result set are identical. A new seed will most likely result in a different sample;
- *Context size*: by entering a number you can determine the number of tokens (words or punctuation marks) Before hit and After hit;
- *Wide View*: the default setting is 'small view'; you can change to Wide View by ticking the checkbox.

Global settings
×

Results per page:

20 results per page ▼

Sample size:

percentage ▼

Sample size ⇅

Seed:

Seed ⇅

Context size:

Context size ⇅

☐ Wide View

Close

Extended search

Like in Simple search, Extended Search allows you to quickly search for specific word forms. The search is performed in the same way as described for Simple Search.

After entering a search term, a spinner briefly appears on the right side of the search bar. Based on the keyed in word, suggestions are given of possible variants of spelling and/or forms from the [GiGaNT-lexicon](#).

Based on the information in this lexicon all spelling variants of the search term found are suggested (see screenshot below).

Search
Explore

Search for ...

Simple
Extended
Advanced
Expert

Word

vlees

Select all

Deselect all

☐ vlees
☐ vleesch
☐ vleesche

☐ vleisch

Limit to Part of Speech

☒ vlees (NOU-C)

☐ Case- and diacritics-sensitive

You can then choose from the presented suggestions or select all at the same time (Select all). To make your search even more targeted, it is also possible to limit the search to certain parts of speech

that were found in the historic component of the GiGaNT-lexicon in connection to the search term. It is also possible to enter a phrase: *wij dronken thee in de salet or wat geeft ons een volkje zonder aenzien.*

In Extended Search it is also possible to search case- and diacritics-sensitive. Note that the default setting for search is case- and diacritics-insensitive. For example, searching for the Word *jan* will result in 301 occurrences. By ticking the box Case- and diacritics-sensitive you will find 14 occurrences of the Word *jan*, but 280 of *Jan* (and 7 of *JAN*). In order to directly find only occurrences of the Word (form) *Jan*, use the search term *Jan* and tick the box Case- and diacritics-sensitive under the search field Word (as shown below).

The screenshot shows the 'Extended' search tab selected. The search field 'Word' contains the text 'Jan'. Below the field are 'Select all' and 'Deselect all' buttons. A list of search results is displayed with checkboxes: ☒ jan, ☒ jans, ☒ ian, ☒ jannen, and ☒ Jan (NOU-P NOU-C). Under the 'Limit to Part of Speech' section, ☐ jan (NOU-C) and ☐ gunnen (VRB) are unchecked, while ☒ Jan (NOU-P NOU-C) is checked. At the bottom, the 'Case- and diacritics-sensitive' checkbox is also checked.

If you know exactly which word you're looking for, you can also – while the wheel is spinning – press Enter directly. The search will then start immediately.

Like in Simple Search, wildcards are supported in Extended Search. (See for a short explanation of wildcards [Simple Search](#))

In the search field Word it is possible to search for different values simultaneously by separating them without spaces by a vertical line, e.g. *god|man|lief* or – with the use of wildcards – *god|aan*|hond*.

For the search field Word it is possible to search for a series of tokens by entering multiple values – including wildcards – separated by a space, e.g. *goede god*, *goede ** or ** god*. It will be obvious that these three searches give different results.

Starting a new search

You can start a new search by pressing the Reset button. By doing so, both the search query and the hits found will disappear. Your search history, however, will remain unchanged.

There are two possibilities to start a search: fill in the desired value and press enter or fill in the desired value and then click the Search button.

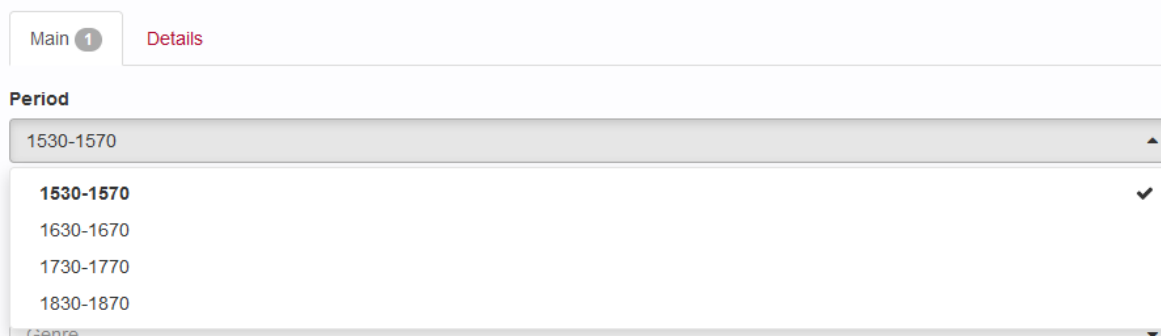
Filter search by

At the right side you will find the option to limit your query to a subset of documents with specific metadata values. You can apply different filters for Main (*Period*, *Region*, *Genre*) and Details (*Author*,

Place, Year). To view the results for all documents, simply leave the attributes in the filtering form empty.

There are two different ways to specify a filter, depending on the field type. Most fields allow you to choose one or more values from a drop-down list, while Year allows you to fill in the value(s) yourself.

Filter search by ...

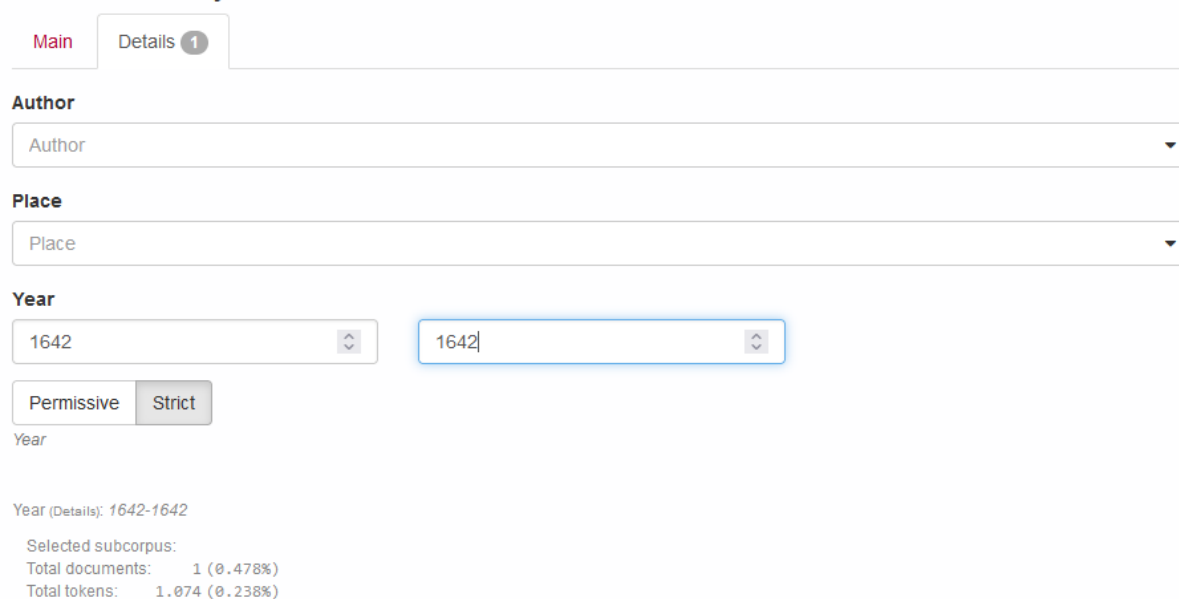


By means of a number at the top of Filter search by, the number of values used to filter on, is displayed as can be seen in the above screenshot.

Filter by year

The documents in this corpus were written or printed in the period between 1530 and 1870. It is either possible to select a preset range of periods under the Main tab (Period) or to select your own range under the Details tab (Year). You can find documents from a specific year by entering the same year in the “from” row as in the “to” row (see screenshot below). If you do not enter a specific year, the entire corpus is searched. If you want to filter by another year or another period, please press the “reset” button.

Filter search by ...



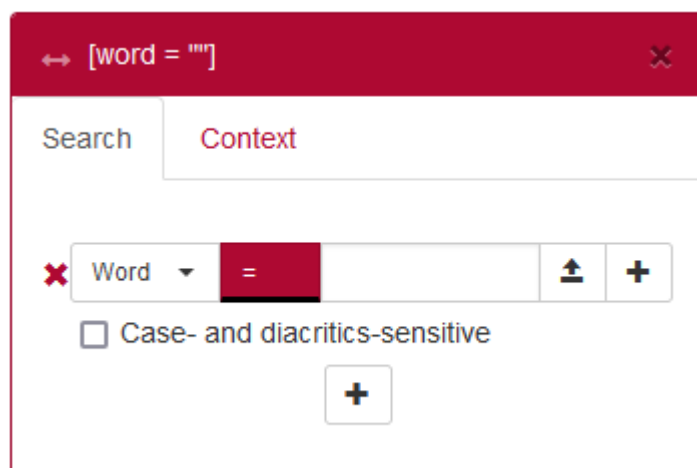
For a detailed description of the metadata, see the section [Metadata categories](#).

Advanced search

The query builder

The basic building block in the query builder is the *token box* (see below). Each box represents a token – usually just a single word – or a simple repetition of tokens; when multiple tokens are used, they are matched in order from left to right.

You can use the query builder to create complex queries without writing CQL (here: Corpus Query Language). Therefore, it is easy to use.



The screenshot shows a 'token box' interface. At the top, a red header bar contains a double-headed arrow icon, the text '[word = ""]', and a close 'X' button. Below the header, there are two tabs: 'Search' (active) and 'Context'. The main area contains a search configuration section. It starts with a red 'X' icon, followed by a dropdown menu set to 'Word', an equals sign '=' in a red box, and an empty text input field. To the right of the input field are two buttons: one with an up/down arrow and another with a plus sign '+'. Below this section is a checkbox labeled 'Case- and diacritics-sensitive' which is currently unchecked. At the bottom center of the box is a large plus sign '+' button.

A token box in the querybuilder has two tabs: Search and Context.

The tab search

The tab search contains a set of attributes a token in the corpus must have to be matched by the query. By clicking the + -button on the right hand side of this token, you can add new attributes (see below). Then enter a value that the attribute must have for the token to be found. The search command Word 'starts with' *ge* and Word 'ends with' *en* for example results in both verbal forms (*gebieden*, *gebeurden*, *gebonden*) and plural nouns (*gecommitteerden*, *gedeputeerden*).

It is only possible to search by word forms. However, you can specify whether that word form should be equal or not equal to the entered search term. You can also specify whether or not a word should begin or end with a particular letter combination.

The CQL query generated to match this token (the *token query*) in the corpus is displayed in the top bar of the box, to help you understand what is happening internally. The following applies to our example:

The screenshot shows a search query builder window with a red header bar containing the query `[word = "ge.*" & word = ".*den"]` and a close button. Below the header, there are two tabs: "Search" and "Context". The "Search" tab is active. It contains two query clauses separated by an "AND" operator. The first clause is "Word" with a dropdown menu showing "star" and the value "ge". The second clause is "Word" with a dropdown menu showing "end" and the value "den". Both clauses have a checkbox for "Case- and diacritics-sensitive" which is unchecked. There are plus and minus buttons for each clause and a plus button at the bottom to add more clauses.

Token attributes

Specifying token attributes is similar to the Extended Search form. Select which attribute a token should have, and enter the value that the attribute must have for the token to be matched. Attributes in the query builder are interpreted as *regular expressions*. Note that this is different from the Extended Search, where token patterns use wildcards.

Going beyond single-attribute token queries, a token box also allows you to combine several attributes and to specify repetition options.

Adding attributes to a token box

Using the +-button, new attributes can be added. Two options exist: *AND* and *OR*.

The *AND* option creates a new attribute restriction that a token must match in addition to the ones which were already there. As an example: suppose we want to match past participles of strong verbs. First, fill in the attribute Word 'starts with' *ge*, then click +, choose *AND*, and choose Word 'ends with' *en*.

This screenshot is identical to the one above, showing the same AND query: `[word = "ge.*" & word = ".*den"]`.

Similarly, creating a new attribute using *OR* will create a token query matching tokens that have the original attribute *or* the new attribute. For instance, enter Word 'starts with' *ge*, add a new attribute with the *OR* option and enter Word 'ends with' *en* to match tokens as *geselschap*, *gesleept* and *bydien*, *afcommen*.

The screenshot shows the search query builder window with the query `[word = "ge.*" | word = ".*en"]`. The operator between the two clauses is "OR". The rest of the interface is the same as the previous screenshots.

Function of the two +-buttons in a token box

The difference between the +-sign on the right of an attribute and the one below it, is that the +-sign on the right keeps the newly added attribute “within a subclause”. This is most easily explained by means of an example.

Suppose we want to look for past participles of weak verbs, i.e. end in an *-d* or a *-t*. If we add the attributes in the order Word ‘starts with’ *ge* AND Word ‘ends with’ *d*, OR Words ‘ends with’ *t* using the +-signs **below** the attributes, as in the left screenshot below, we get the token query [(word = "ge.*" & word = ".*d") | word = ".*t"]. This will also match forms such as *stadt*, *niet*, so this is not what we were after.

If, on the other hand, we add OR Word ‘ends with’ *t* with the +-sign to the **right** of the attribute Word ‘ends with’ *d*, it will be inserted in a subclause, thus resulting in the correct query [word = "ge.*" & (word = ".*d" | word = ".*t")], as shown in the right screenshot below.

The left screenshot shows a search interface with a query bar containing `[(word = "ge.*" & word = ".*d") | word = ".*t"]`. Below the query bar, there are three attribute boxes. The first box is for 'Word' starting with 'ge'. The second box is for 'Word' ending with 'd', and it is connected to the first box by an 'AND' operator. The third box is for 'Word' ending with 't', and it is connected to the second box by an 'OR' operator. The right screenshot shows a similar interface, but the third box is for 'Word' ending with 't', and it is connected to the second box by an 'OR' operator. The query bar now contains `[word = "ge.*" & (word = ".*d" | word = ".*t")]`.

The tab Context

The tab options specifies the contextual properties, such as whether the token occurs at the end of a sentence, and the repetition pattern:

The screenshot shows the 'Context' tab of the search interface. It has a query bar with `[word = ""]`. Below the query bar, there are four checkboxes: 'Optional', 'Begin of sentence', and 'End of sentence'. At the bottom, there is a repetition pattern section with a 'repeats' field set to '1', a 'to' field set to '1', and a 'times' field.

Managing sequences of token boxes

There are three ways to manage the sequence and the number of token boxes:

- *Rearrange* a token by clicking and dragging the little arrow handle in the top-left corner simultaneously (1).
- *Delete* a token by clicking the **x** in the top-right corner (2).
- *Create a new token box* by clicking the **+**-button next to the upper right corner of the utmost right token box (3).



Uploading value lists in the query builder

It is also possible to upload a list of values, separated by a white space. To do so, click the upload button (with the arrow pointing upwards) and select a text file. Tokens will then be matched for any of the values from the file.

Note that this function only works for *.txt-files. If you are using a text editor like Word, you have to save your file as a *.txt file or you can copy and paste the values into a *.txt file first.

After uploading a file, the text can be edited by clicking the yellow marked file name in the text field. Editing the text is temporary and will not modify your original file.

To remove an uploaded file and go back to typing a value, click on the cross (x) next to the yellow text box. Another possibility to clear the uploaded values is by clicking the yellow marked text field and then pressing the Clear button on the bottom left corner of the Edit box. Using the Reset button will start a complete new search.

Copy to CQL editor

You can use the query builder to create complex queries without writing CQL. Any time a query is created in the querybuilder, it can be copied to the CQL editor, where you can further edit the query by hand. This will take you automatically to the Expert Search screen, after which you can start the search or adjust the query if desired.

Copy to CQL editor

Expert search

The Corpus Query Language (CQL) editor allows you to type your own CQL query, to import a previously downloaded query and to upload a tab separated list of values to substitute for gap values (see below for further explanation).

CQL queries are expressions built up with the help of a few sequence operators and brackets from basic blocks enclosed by square brackets, in each of which one or more token attributes are specified.

In CQL, spaces only affect a search if they are included in quotes. Whether the search command is `[word="schip"]` or `[ word = "schip" ]` (or just “schip”) does not make any difference to the result. However, there is a difference between the queries `[word="schip"]` and `[word=" schip"]`. The first search results in exactly 123 hits, but the second one in zero!

Some examples:

- Simple: `[word="hand"]`, e.g. the attribute word matches the regular expression *hand*; `[word!="hand"]`, e.g. the attribute word does **not** match the regular expression *hand*; `[word="*.man"]` matches all words ending with *man*, including *man* itself. (Note that `[word="*man"]` will not give any results, because in Expert Search an asterisk is not a wildcard but a repetition operator.)
- Combination of attributes (combining operators are &, |, !), e.g. `[word="hoop|geloof|liefde"]` matches either the word *geloof*, the word *hoop* or the word *liefde*.
- The empty `[]` matches any token, e.g. `[word="man"][]{}3[word="ik"]` matches a sequence of *man* followed by *ik* with three arbitrary tokens in between.
- Operators |, & and parentheses () and the repetition operators (+, *, ? and {}) can be used to build complex sequence queries. Example: `"laatste" "tijd" | "eerste" "maal"`, matching any sequence of *laatste tijd* or *eerste maal*.

This short list does not cover all CQL features. For more detailed information on how to write CQL, please consult the short [Appendix: Corpus Query Language](#), which contains further pointers.

Copy to query builder

When the query is relatively simple - like `[word="schip"][word="den"]` - it can also be imported into the querybuilder using the *Copy to query builder* button. This will take you automatically to the Advanced Search screen, after which you can start the search or adjust the query if desired.

A message will be displayed next to the button if the query couldn't be parsed.

Import query

If you have entered a search query, you can find it back by clicking the History button. On the right hand side you can select Download as file in the drop-down menu (default value is Search) and save the file. (For a more elaborate description of the History button see [Simple Search](#).)

Previously saved queries can be used again by uploading them through the Import query button.

Gap filling

Use this button to upload a Tab Separated Values (TSV) file, which is a simple text format for storing data in a tabular structure. Each record in the table is one line of the text file. Each field value of a record is separated from the next by a tab character. It is also possible to upload a plain text file (.txt) that has the same properties.

A *.tsv file or a comparable *.txt file enables you to complete a query with marked gaps.

If, for instance, you are interested in the distribution of adjectives you can create this query in the Corpus Query Language field:

```
[word="@@"][][word="@@"]
```

By clicking Gap-filling you can upload a file with a tab-separated list of values from your computer to substitute them for the gap values, i.e. the at signs (@@) in your query. After the upload your values will appear in a separate box:

Corpus Query Language: ⓘ

```
[word="@@"][][word="@@"]
```

Copy to query builder
Import query
Gap-filling ✕

de	man
een	vrouw
het	huis

The values in the first column – *de*, *een*, *het* – will be entered at the position of the first gap (@@) and the values in the second column – *man*, *vrouw*, *huis* – at the position of the second gap. With these values, gap-filling yields the following results:

Before Hit	Hit	After Hit
KB Pfl18264 (Pamphlets), Amsterdam, 1750 ... en tot 14 toe gedeserteert,	de tien man	de couragie behoudende, zyn boven...
KB Pfl17845 (Pamphlets), Amsterdam, 1748 ... wanneer dat quam te excedeeren	de 6000 Man	, zoo zullen door den Capitein...
SAB Notulen gemeenteraad 1830 (Administrative), Brugge, 1830 ... bouwen eener Kaai muur, langs	het aanbelindende huis het zelve huis	, toebehoorende aan den heer Jacobus... te ontgaan; het beloop van...
1549 Jan De Pottre (Ego-documents), Brussel, 1549-1570 ... veel peerde cavaliers spaennaerts die	de aermen man	zeer grooten oever laest deden...
1763 DeNeufville 01-07 (Ego-documents), Amsterdam, 1763 ... in het parloir Bragt dat	een Charmante Vrouw	was (zynde ook van geboorte...
1829 V. van Schaeybroek (Ego-documents), Kessel-Lo, 1860 ... den kant van het Evangelie,	een Lieven vrouw	beld met eenen slegten troon...
KB Pfl17781 (Pamphlets), Middelburg, 1747 ... luy is: en toe om	de moye man	te weesen, was y aast...

This mimics the functionality to upload a list of values in the Extended Search and Advanced Search interfaces.

Please note that for this to work, you do need to enter @@ in the field where you want the substitution to take place. An empty field ([]) will match any term.

Viewing results

Results can be viewed in two ways: Per hit (hit is defined as one token or a group of tokens that matched the query), or Per document (each document listed contains at least one hit).

Per Hit view

Click a hit – i.e. a line with the bold word(s) in the column Hit – to display the properties and values of the hit (in the following example **het aanbelindende huis**). Click the hit again to close.

Document id: ADM-1850-VL-3	Before Hit	Hit	After Hit
SAB Notulen gemeenteraad 1830 (Administrative), Brugge, 1830			
...bouwen eener Kaai muur, langs het aanbelindende huis , toebehoorende aan den heer Jacobus...			
...verfraaijing van de stad zal medewerken, en na den toegang of overgang dier brug meer veiligheid toebrengen; Overwegende dat door de uitvoering van dat ontwerp , de stad in de noodzakelijkheid zoude worden gesteld, om eenen onkost van ongeveer twee duizend guldens te doen, voor het bouwen eener Kaai muur, langs het aanbelindende huis , toebehoorende aan den heer Jacobus Berger , welk huis aan den hoek van de Molenmeersch-sstraat A. 6, no. 1, ten noordooste van het hoofd der Brug, al den kant van de lange straat , gelegen is, welken onkost men zoude kunnen vermijden, indien de stad gemelde huis zoude aankooopen, behoudens het zelve...			
Property	value		
Word	het	aanbelindende	huis

Hit rows are always preceded by a row containing the document title in which those hits occurred, in this case *SAB Notulen gemeenteraad 1830 (Administrative), Brugge, 1830*. The document titles can be toggled on or off by using the Hide Titles (or Show Titles when titles are hidden) button at the bottom of the page. If you hover the mouse over the title, the identification number of the document appears, in this case: Document id: ADM-1850-VL-3.

Sorting results

Click on any of the column headings to sort the hits on Words within that column, clicking again inverts the sorting.

You can also sort the results by means of the drop-down menu at the bottom of the page (Sort by...), which offers you the possibility to sort by various attributes such as Hit, Before hit, After hit, Main, Details, Metadata.

Grouping results

It is possible to group the results by clicking on the button Group Results, after which the following menu appears:

Group Results

+ Annotation

+ Metadata

Click on Annotation or Metadata to define grouping criteria.

Close

Group

Results can be grouped by Annotation and by Metadata.

By clicking +Annotation you can group by the first word, by all words or by specific words, whether before the hit, within the hit or after the hit, and based on the annotation Word. Clicking +Metadata allows you to group by metadata assigned to the document (Main, Details and Metadata).

The default grouping is grouping all words within the hit using annotation Word. By clicking the Case sensitive box it is possible to distinguish between case sensitive and case insensitive.

The example below is grouped by the first word before the hit. The example dynamically updates when the grouping options are changed.

Group Results
+ Annotation
+ Metadata

first Word before hit ✕

I want to group on
the first word ▾
before the hit ▾
using annotation
Word ▾

Case sensitive: ☐

daer naestuolgende noch vijftich **carolus** gulden daeraff ende tsurplus Inde ganghinghe
daer naestuolgende noch vijftich carolus gulden daeraff ende tsurplus Inde ganghinghe

Clear
Group

Click a group to show or hide hits within that group, as shown below. Click once more on the group to close it again. If more than twenty hits are found in a document, you can make them appear by clicking on Load more concordances.

Group	#hits in group	Relative frequency (hits)																																																															
karolus	21	0.00465%																																																															
View detailed concordances Load more concordances <table> <thead> <tr> <th>Before Hit</th><th>Hit</th><th>After Hit</th></tr> </thead> <tbody> <tr><td>...Om ende midts vyffthien karolus</td><td>gulden</td><td>eens metter naest vallender Rente...</td></tr> <tr><td>...van tveysch ende een karolus</td><td>gulden</td><td>zoe menichwerff yemant dair op...</td></tr> <tr><td>...opde verbeurte van drie karolus</td><td>gulden</td><td>Item dat nyemant duyebreet van...</td></tr> <tr><td>...zal verbueren een halue karolus</td><td>gulden</td><td>tot profyt vanden voorseide bewaerder...</td></tr> <tr><td>...van elcke Reyse drie karolus</td><td>gulden</td><td>ende dair toe die correctie...</td></tr> <tr><td>...die zal gheuen zes karolus</td><td>gulden,</td><td>ten zy dat hy hylckte...</td></tr> <tr><td>...opte verbuerte van twe karolus</td><td>gulden,</td><td>Item dat alle hudecoopers poorters...</td></tr> <tr><td>...de boeten van twe karolus</td><td>gulden</td><td>elcke Reyse yemant hier op...</td></tr> <tr><td>...op pene van twe karolus</td><td>gulden</td><td>elcke Reyse te verbueren Item...</td></tr> <tr><td>...opde pene van twee karolus</td><td>gulden</td><td>als voeren. Alle welcke boeten...</td></tr> <tr><td>...op peene van twaelf karolus</td><td>gulden</td><td>Item dat de coopers vant...</td></tr> <tr><td>...somme noch geuen twelf karolus</td><td>gulden</td><td>eens, ende dat zo verre...</td></tr> <tr><td>...ende daertoe van een karolus</td><td>gulden</td><td>op elck vat ofte tonnebiere...</td></tr> <tr><td>...op peene van zes karolus</td><td>gulden</td><td>te verbueren. zoe dick ende...</td></tr> <tr><td>...opde boete van drie karolus</td><td>gulden</td><td>elcke Reyse tappliceren als bouen...</td></tr> <tr><td>...te hebben, van zes karolus</td><td>gulden</td><td>tot des heeren proufyt ende...</td></tr> <tr><td>...somme van zes hondert karolus</td><td>gulden,</td><td>waer an de voorscreuen coopluden...</td></tr> <tr><td>...aen te verbueren zes karolus</td><td>gulden,</td><td>tappliceren als bouen, ende daer...</td></tr> <tr><td>...bouen te verbueren zes karolus</td><td>gulden,</td><td>ende den geenen diet voor...</td></tr> <tr><td>...op peyne van zes karolus</td><td>gulden,</td><td>elcke reyse te verbueren ende...</td></tr> </tbody> </table> View detailed concordances Load more concordances			Before Hit	Hit	After Hit	...Om ende midts vyffthien karolus	gulden	eens metter naest vallender Rente...	...van tveysch ende een karolus	gulden	zoe menichwerff yemant dair op...	...opde verbeurte van drie karolus	gulden	Item dat nyemant duyebreet van...	...zal verbueren een halue karolus	gulden	tot profyt vanden voorseide bewaerder...	...van elcke Reyse drie karolus	gulden	ende dair toe die correctie...	...die zal gheuen zes karolus	gulden,	ten zy dat hy hylckte...	...opte verbuerte van twe karolus	gulden,	Item dat alle hudecoopers poorters...	...de boeten van twe karolus	gulden	elcke Reyse yemant hier op...	...op pene van twe karolus	gulden	elcke Reyse te verbueren Item...	...opde pene van twee karolus	gulden	als voeren. Alle welcke boeten...	...op peene van twaelf karolus	gulden	Item dat de coopers vant...	...somme noch geuen twelf karolus	gulden	eens, ende dat zo verre...	...ende daertoe van een karolus	gulden	op elck vat ofte tonnebiere...	...op peene van zes karolus	gulden	te verbueren. zoe dick ende...	...opde boete van drie karolus	gulden	elcke Reyse tappliceren als bouen...	...te hebben, van zes karolus	gulden	tot des heeren proufyt ende...	...somme van zes hondert karolus	gulden,	waer an de voorscreuen coopluden...	...aen te verbueren zes karolus	gulden,	tappliceren als bouen, ende daer...	...bouen te verbueren zes karolus	gulden,	ende den geenen diet voor...	...op peyne van zes karolus	gulden,	elcke reyse te verbueren ende...
Before Hit	Hit	After Hit																																																															
...Om ende midts vyffthien karolus	gulden	eens metter naest vallender Rente...																																																															
...van tveysch ende een karolus	gulden	zoe menichwerff yemant dair op...																																																															
...opde verbeurte van drie karolus	gulden	Item dat nyemant duyebreet van...																																																															
...zal verbueren een halue karolus	gulden	tot profyt vanden voorseide bewaerder...																																																															
...van elcke Reyse drie karolus	gulden	ende dair toe die correctie...																																																															
...die zal gheuen zes karolus	gulden,	ten zy dat hy hylckte...																																																															
...opte verbuerte van twe karolus	gulden,	Item dat alle hudecoopers poorters...																																																															
...de boeten van twe karolus	gulden	elcke Reyse yemant hier op...																																																															
...op pene van twe karolus	gulden	elcke Reyse te verbueren Item...																																																															
...opde pene van twee karolus	gulden	als voeren. Alle welcke boeten...																																																															
...op peene van twaelf karolus	gulden	Item dat de coopers vant...																																																															
...somme noch geuen twelf karolus	gulden	eens, ende dat zo verre...																																																															
...ende daertoe van een karolus	gulden	op elck vat ofte tonnebiere...																																																															
...op peene van zes karolus	gulden	te verbueren. zoe dick ende...																																																															
...opde boete van drie karolus	gulden	elcke Reyse tappliceren als bouen...																																																															
...te hebben, van zes karolus	gulden	tot des heeren proufyt ende...																																																															
...somme van zes hondert karolus	gulden,	waer an de voorscreuen coopluden...																																																															
...aen te verbueren zes karolus	gulden,	tappliceren als bouen, ende daer...																																																															
...bouen te verbueren zes karolus	gulden,	ende den geenen diet voor...																																																															
...op peyne van zes karolus	gulden,	elcke reyse te verbueren ende...																																																															
carolus	19	0.00421%																																																															
duysent	9	0.00199%																																																															

Click on View detailed concordances to go back to the normal hits view to see more detailed information for the hits in this group. The button Go back to grouped results brings you back to the list of groups.

Per Document view

Sorting results

Click on any of the column headings to sort the documents by (document) name, period, region, genre or hits within that column, clicking again inverts the sorting.

You can also sort the results by means of the drop-down menu at the bottom of the page (Sort by...), which offers you the possibility to sort by various attributes such as Hit (Documents), Period (Main) and Author (Details).

Grouping results

Results Per Document can be grouped by metadata assigned to the document (Main, Details and Metadata). The example below shows all documents in which the Word *man* occurs grouped by period.

Results for: [word="man"] within all documents

Per Hit Per Document

Documents / Grouped by Document Period

Total documents: 65 (31.1%)
Total groups: 4
Search time: 0.001s

Group Results + Metadata

Document Period ✕

Select the metadata to group on.

Group by Period ▾

Case sensitive: ☐

Clear Group

« 1 » table docs

Group	#docs in group	Relative frequency (docs)
1630-1670	22	33.8%
1730-1770	17	26.2%
1530-1570	15	23.1%
1830-1870	11	16.9%

Exporting results

The search results – both Per Hit as Per Document – can be exported by using the Export or the Export for Excel button at the bottom right of the page. The first button transfers the search results – including all metadata – to a Comma-Separated Values-file. These CSV-files consist only of text data, which makes it easy to implement (read and/or write) them into a spreadsheet or database program. The second button offers the possibility to export the results – including all metadata – to a CSV-file for use with Excel.

Grouped results can be exported in the same way. However, if you would like to have the metadata with each concordance of a group, you must first click on the red bar of a specific group and then on View detailed concordances. The results you then see can be exported by the use of the Export buttons. This operation must be carried out for each individual group you wish to export.

Information about a document

Click on a document title to open the document in a new window: the Content window.

Content

Hits from the current query will be highlighted in bold in the opened document. You can navigate from one hit to another by using the arrows at the Hits button:

SAL Schepenbank Leuven 1553-1568 (Administrative), Leuven, 1553-1568

[schepenjaar 1553, Folio 18R]

Nae dien den Raide der **stadt** van Louenen
by supplicatien te kennen gegeuen Is van wegen
Jans van Rillaer schilder van deser **stadt**
dat hy heeft drie Jaeren te wetene Inden
Jaeren vyftich eenenvyftich ende tweenvyftich
gewracht Inder **stadt** halle telcken
zesse weken voer Louen kermisse zonder
ophouden om te ouersiene alle datter
nootelyck zyn mochte ende tselue alsoe
te versiene ende te Repareren datter gheen

Hit « 1/18 » »

Metadata

In the metadata tab, all metadata properties of the document are displayed. They provide information about Main (*Period*, *Region*, *Genre*) and Details (*Author*, *Place*, *Year*) and also Word count and Document Length.

Statistics

The Statistics tab shows several document statistics: the number of Tokens, Types and the Type/Token ratio. It is possible to print or to download these statistics via the menu symbol right of the title Vocabulary Growth.

Exploring the corpus

The Explore tab has three subdivisions: Documents, N-grams and Statistics.

Documents

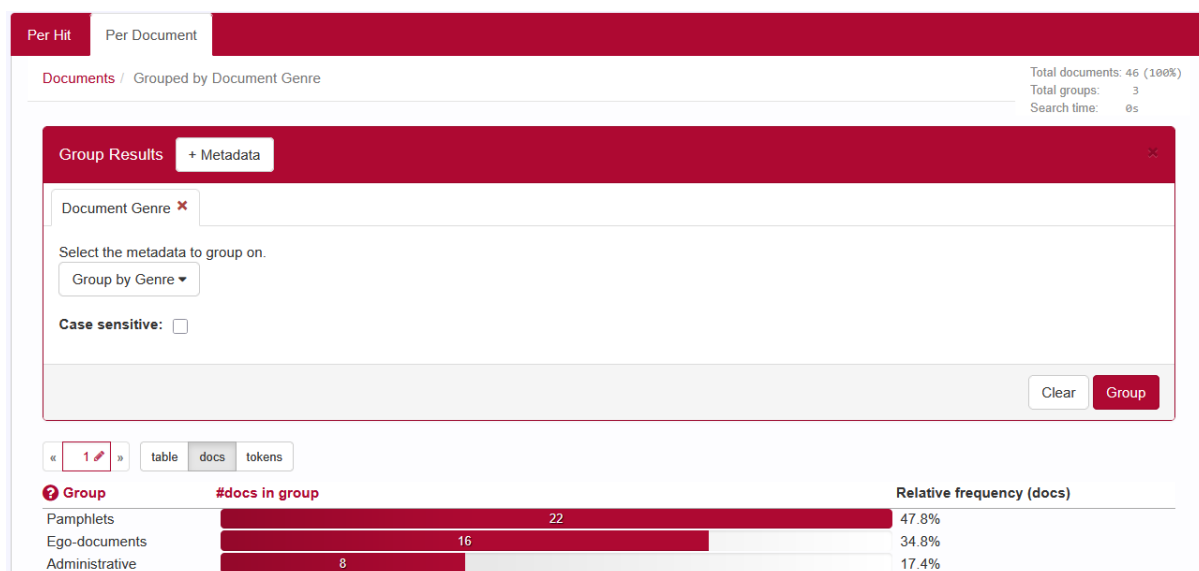
This subtab allows you to investigate the documents. It consists of two drop-down menus to specify the grouping of the metadata and to specify the way the groups are to be shown.

A simple example: suppose we want to know the genres of the documents from Zeeland.

- In the Group documents by metadata drop-down menu, choose Group by Genre
- In Show groups as, select *Docs*
- In the metadata search form (Filter search by), select in Region *Zeeland*
- Press Search

The 'Explore' interface has a top navigation bar with 'Search' and 'Explore' tabs. Below the navigation bar, there are three tabs: 'Documents', 'N-grams', and 'Statistics'. The 'Documents' tab is active. Under 'Explore ...', there are two dropdown menus: 'Group documents by metadata' (set to 'Group by Genre') and 'Show groups as' (set to 'Docs'). Below this, there is a 'Filter search by ...' section with two tabs: 'Main' (with a count of 1) and 'Details'. Under 'Filter search by ...', there are two dropdown menus: 'Period' (set to 'Period') and 'Region' (set to 'Zeeland').

You will get the following result:



N-grams

An *N-gram* is a sequence of *N* items. This option will list the frequency of different N-grams in a (sub-)corpus.

Options

- *N-gram size*: the length of the sequence (a number from 1 to 5; default setting is 5)
- *N-gram-type*: a sequence of five consecutive Words (i.e. word forms). If you do not specify the search term a series of five arbitrary words will be searched for.
- It is also possible to restrict to, for instance, 5-grams with some slots already specified, as is shown in the following example. After entering a search term, a spinner briefly appears on the right side of the search bar. Based on the keyed in word, suggestions are given of possible variants of spelling and/or form from the GiGaNT-lexicon and of parts of speech. By clicking on 'Select all' all forms belonging to a GiGaNT lemma are added.

- By using the Filter search by, you can create a subcorpus within the HCD for specific metadata.

Example

Search Explore

Explore ...

Documents N-grams Statistics

N-gram size: 5

N-gram type: Word

Word Word Word Word Word

dees|deeze|deezen|dese|de: man|manne|mannen|mans Word Word Word

Select all Deselect all Select all Deselect all

☒ dees ☒ man

☒ deeze ☒ manne

☒ deezen ☒ mannen

☒ dese ☒ mans

☒ desen ☐ Limit to Part of Speech

☒ deser ☒ man (NOU-C)

☒ deses ☒ manna (NOU-C)

☒ deuse

☒ deze

☒ dezen

☒ dezer

Limit to Part of Speech

☒ deze (PD)

Within all the documents of the HCD, you will find 4 occurrences of this so-called 5-gram.

Per Hit Per Document

Hits / Grouped by Word within hit

Total hits: 4 (0.000886%)
Total groups: 4
Search time: 0.02s

Group Results + Annotation + Metadata

Word within hit ✕

I want to group on all words within the hit using annotation Word

Case sensitive: ☐

tot Colen toe Ende reden | **desen man dach wt mastricht** | ij milen tot Gulpen berch
toi. Colen toe Ende reden | desen man dach wt mastricht | . milen toi. Gulpen berch

Clear Group

« 1 » table hits

Group	#hits in group	Relative frequency (hits)
dese man daar weer quam	1	0.000222%
deser Mannen soo veel noch	1	0.000222%
dese man den Ommegank gemaakt	1	0.000222%
desen man dach wt mastricht	1	0.000222%

Statistics (frequency lists)

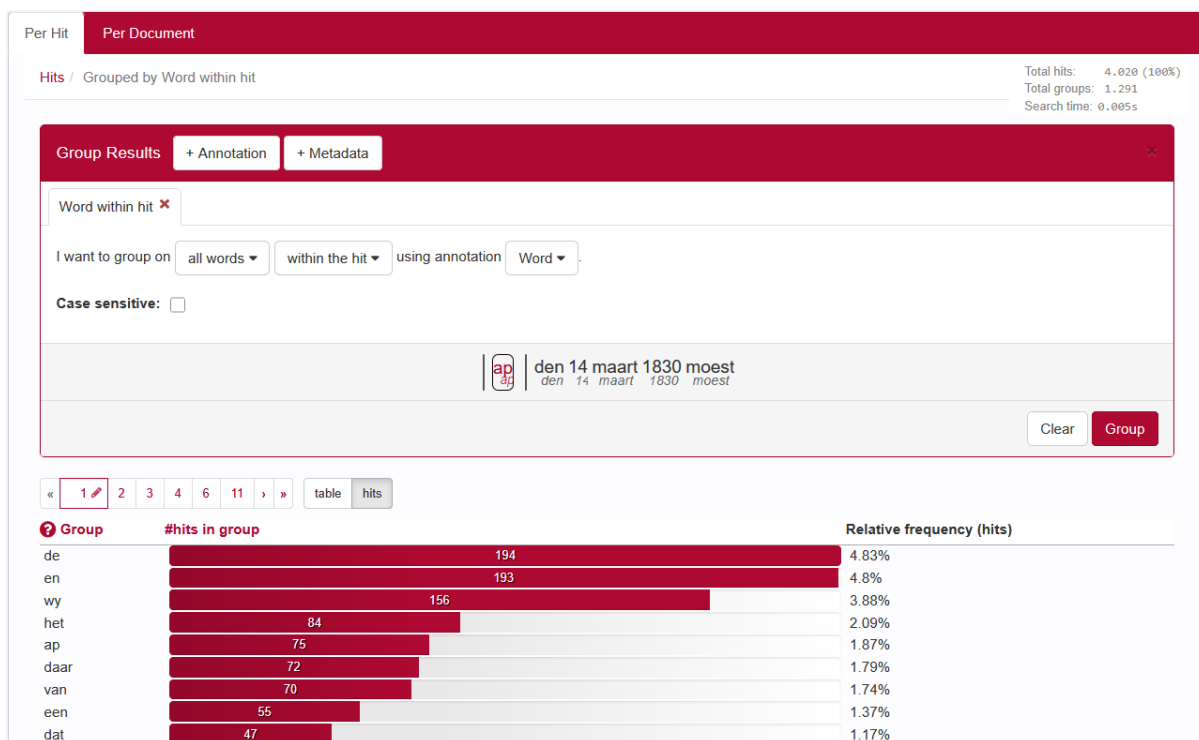
Here, you can produce frequency lists for the corpus. It is rather similar to the previous option, but restricted to 1-grams.

Options

- *Frequency list type*: in this corpus, it is only possible to create frequency lists of Words (i.e. word forms)
- By using the Filter search by, you can create a subcorpus within the HCD for specific metadata.

Example

It is possible to determine the use of the ten most frequently used words in texts from Alkmaar in the HCD by searching for Frequency list type Word and by filtering search by Place: *Alkmaar*. This results in:



Appendix: Corpus Query Language

BlackLab supports Corpus Query Language, a full-featured query language introduced by the IMS Corpus WorkBench (CWB) and also supported by the Lexicom Sketch Engine. It is a standard and powerful way of searching corpus.

The basics of Corpus Query Language is the same in all three projects, but there are a few minor differences in some of the more advanced features, as well as some features that are exclusive to some projects. For most queries however, this will not be an issue.

This page will introduce the query language and show all features that BlackLab supports. If you want to learn even more about CQL, see [CWB CQP Query Language Tutorial](#) and [Sketch Engine Corpus Query Language](#).

CQL support

For those who already know CQL, here's a quick overview of the extent of BlackLab's support for this query language. If there is a feature we don't support, yet is important to you, please let us know. If it's quick to add, we may be able to help you out.

Supported features

BlackLab currently supports (arguably) most of the important features of Corpus Query Language:

- Matching on token annotations (also called properties or attributes), using regular expressions and =, !=, !. Example: [word="bank"] (or just "bank")
- Case/accent-sensitive matching. Note that, unlike in CWB, case-INsensitive matching is currently the default. To explicitly match case/accent-insensitivity, use "(?i)...". Example: "(?i)Mr\." "(?i)Banks"
- Combining criteria using &, | and !. Parentheses can also be used for grouping. Example: [lemma="bank" & pos="V"]
- Match-all pattern [] matches any token. Example: "a" [] "day"
- Regular expression operators +, *, ?, {n}, {n,m} at the token level. Example: [pos="AA"]+
- Sequences of token constraints. Example: [pos="AA"] "cow"
- Operators |, & and parentheses can be used to build complex sequence queries. Example: "happy" "dog" | "sad" cat"
- Querying with tag positions using e.g. <s> (start of sentence), </s> (end of sentence), <s/> (whole sentence) or <s> ... </s> (equivalent to <s/> containing ...). Example: <s> "The" . XML attribute values may be used as well, e.g. <ne type="PERS"/> ("named entities that are persons").
- Using within and containing operators to find hits inside another set of hits. Example: "you" "are" within <s/>
- Using an anchor to capture a token position. Example: "big" A:[]. Captured matches can be used in global constraints (see next item) or processed separately later (using the Java interface; capture information is not yet returned by BlackLab Server). Note that BlackLab can actually capture entire groups of tokens as well, similarly to regular expression engines.
- Global constraints on captured tokens, such as requiring them to contain the same word. Example: "big" A:[] "or" "small" B:[] :: A.word = B.word

See below for features not in this list that may be added soon, and let us know if you want a particular

feature to be added.

Differences from CWB

BlackLab's CQL syntax and behaviour differs in a few small ways from CWBs. In future, we'll aim towards greater compliance with CWB's de-facto standard (with some extra features and conveniences).

For now, here's what you should know:

- Case-insensitive search is currently the default in BlackLab, although you can change this if you wish. CWB and Sketch Engine use case-sensitive search as the default. We may change our default in a future major version.
If you want to switch case-/diacritics-sensitivity, use "(?-i).." (case-sensitive) or "(?i).." (case-insensitive, usually the default). CWBs %cd flags for setting case/diacritics-sensitivity are not (yet) supported, but will be added.
- If you want to match a string literally, not as a regular expression, use backslash escaping: "e\.g\.". %l for literal matching is not yet supported, but will be added.
- BlackLab supports result set manipulation such as: sorting (including on specific context words), grouping/frequency distribution, subsets, sampling, setting context size, etc. However, these are supported through the REST and Java APIs, not through a command interface like in CWB. See [BlackLab Server overview](#).
- Querying XML elements and attributes looks natural in BlackLab: <s/> means "sentences", <s> means "starts of sentences", <s type='A'> means "sentence tags with a type attribute with value A". This natural syntax differs from CWBs in some places, however, particularly when matching XML attributes. While we believe our syntax is the superior one, we may add support for the CWB syntax as an alternative.
We only support literal matching of XML attributes at the moment, but this will be expanded to full regex matching.
- In global constraints (expressions occurring after ::), only literal matching (no regex matching) is currently supported. Regex matching will be added soon. For now, instead of A:[] "dog" :: A.word = "happy|sad", use "happy|sad" "dog".
- To expand your query to return whole sentences, use <s/> containing (...). We don't yet support CWBs expand to, expand left to, etc., but may add this in the future.
- The implication operator -> is currently only supported in global constraints (expressions after the :: operator), not in regular token constraints. We may add this if there's demand for it.
- We don't support the @ anchor and corresponding target label; use a named anchor instead. If someone makes a good case for it, we will consider adding this feature.
- backreferences to anchors only work in global constraints, so this doesn't work: A:[] [] [word = A.word]. Instead, use something like: A:[] [] B:[] :: A.word = B.word. We hope to add support for these in the near future, but our matching approach may not allow full support for this in all cases.

(Currently) unsupported features

The following features are not (yet) supported:

- intersection, union and difference operators. These three operators will be added in the future.
For now, the first two can be achieved using & and | at the sequence level, e.g. "double" [] & []

"trouble" to match the intersection of these queries, i.e. "double trouble" and "happy" "dog" | "sad" "cat" to match the union of "happy dog" and "sad cat".

- `_` meaning "the current token" in token constraints. We will add this soon.
- `lbound`, `rbound` functions to get the edge of a region. We will probably add these.
- `distance`, `distabs` functions and `match`, `matchend` anchor points (sometimes used in global constraints). We will see about adding these.
- using an XML element name to mean 'token is contained within', like `[(pos = "N") & !np]` meaning "noun NOT inside in an tag". We will see about adding these.
- a number of less well-known features. If people ask, we will consider adding them.

Using Corpus Query Language

Matching tokens

Corpus Query Language is a way to specify a "pattern" of tokens (i.e. words) you're looking for. A simple pattern is this one:

```
[word="man"]
```

This simply searches for all occurrences of the word "man". If your corpus includes the per-word properties `lemma` (i.e. headword) and `pos` (part-of-speech, i.e. noun, verb, etc.), you can query those as well. For example, to find a form of word "search" used as a noun, use this query:

```
[lemma="search" & pos="NOU-C"]
```

This query would match "search" and "searches" where used as a noun. (Of course, your data may contain slightly different part-of-speech tags.)

The first query could be written even simpler without brackets, because "word" is the default property:

```
"man"
```

You can use the "does not equal" operator (`!=`) to search for all words except nouns:

```
[pos != "NOU-C"]
```

The strings between quotes can also contain wildcards, of sorts. To be precise, they are [regular expressions](#), which provide a flexible way of matching strings of text. For example, to find "man" or "woman", use:

```
"(wo)?man"
```

And to find lemmata starting with "under", use:

```
[lemma="under.*"]
```

Explaining regular expression syntax is beyond the scope of this document, but for a complete overview, see [here](#).

Sequences

Corpus Query Language allows you to search for sequences of words as well (i.e. phrase searches, but with many more possibilities). To search for the phrase "the tall man", use this query:

```
"the" "tall" "man"
```

It might seem a bit clunky to separately quote each word, but this allows us the flexibility to specify exactly what kinds of words we're looking for. For example, if you want to know all single adjectives used with man (not just "tall"), use this:

```
"an? | the" [pos="AA"] "man"
```

This would also match "a wise man", "an important man", "the foolish man", etc.

Regular expression operators on tokens

Corpus Query Language really starts to shine when you use the regular expression operators on whole tokens as well. If we want to see not just single adjectives applied to "man", but multiple as well:

```
"an? | the" [pos="AA"] + "man"
```

This query matches "a little green man", for example. The plus sign after [pos="AA"] says that the preceding part should occur one or more times (similarly, * means "zero or more times", and ? means "zero or one time").

If you only want matches with two or three adjectives, you can specify that too:

```
"an? | the" [pos="AA"] { 2 , 3 } "man"
```

Or, for two or more adjectives:

```
"an? | the" [pos="AA"] { 2 , } "man"
```

You can group sequences of tokens with parentheses and apply operators to the whole group as well. To search for a sequence of nouns, each optionally preceded by an article:

```
("an? | the"? [pos="NOU-C"] ) +
```

This would, for example, match the well-known palindrome "a man, a plan, a canal: Panama!"

Punctuation

In BlackLab, punctuation tends to not be indexed as a separate token, but as a property of a word token - CWB and Sketch Engine on the other hand tend to index punctuation as a separate token instead. You certainly could choose to index punctuation as a separate token in BlackLab, by the way -- it's just not commonly done. Both approaches have their advantages and disadvantages, and of course the choice affects how you write your queries.

It is possible to search for punctuation marks. E.g. to find occurrences of the word "want" preceded by a comma use the following query:

```
[punctBefore=", " & word="want"]
```

To find occurrences of the lemma "krant" that are followed by an exclamation mark, use:

```
[lemma="krant" & punctAfter="!"]
```

Some punctuation marks have a special function in regular expressions and therefore must be preceded by a backslash (\) when used in queries. For instance, to search for a period (.) after the word **"geweest"**, use:

```
[word="sentence" & punctAfter="\."]
```

Case- and diacritics-sensitivity

CWB and Sketch Engine both default to (case- and diacritics-)sensitive search. That is, they exactly match upper- and lowercase letters in your query, plus any accented letters in the query as well. BlackLab, on the contrary, defaults to *IN*sensitive search (although this default can be changed if you like). To match a pattern sensitively, prefix it with "(?-i)":

```
"(?-i) Panama"
```

If you've changed the default search to sensitive, but you wish to match a pattern in your query insensitively, prefix it with "(?i)":

```
[pos="( ?i) NOU-C"]
```

Although BlackLab is capable of setting case- and diacritics-sensitivity separately, it is not yet possible from Corpus Query Language. We may add this capability if requested.

Matching XML elements

Corpus Query Language allows you to find text in relation to XML elements that occur in it. For example, if your data contains sentence tags, you could look for sentences starting with "the":

```
<s>"the"
```

Similarly, to find sentences ending in "that", you would use:

```
"that"</s>
```

You can also search for words occurring inside a specific element. Say you've run named entity recognition on your data and all person names are surrounded with <person>...</person> tags. To find the word "baker" as part of a person's name, use:

```
"baker" within <person/>
```

Note the forward slash at the end of the tag. This way of referring to the element means "the whole element". Compare this to <person>, which means "the element's open tag", and </person>, which means "the element's close tag".

The above query will just match the word "baker" as part of a person's name. But you're likely more interested in the entire name that contains the word "baker". So, to find those full names, use:

```
<person/> containing "baker"
```

Or, if you simply want to find all persons, use:

```
<person/>
```

As you can see, the XML element reference is just another query that yields a number of matches. So as you might have guessed, you can use "within" and "containing" with any other query as well. For example:

```
([pos="AA"]+ containing "tall") "man"
```

will find adjectives applied to man, where one of those adjectives is "tall".

Labeling tokens, capturing groups

Just like in regular expressions, it is possible to "capture" part of the match for your query in a "group".

CWB and Sketch Engine offer similar functionality, but instead of capturing part of the query, they label a single token. BlackLab's functionality is very similar but can capture a number of tokens as well. For example:

```
"an?|the" Adjectives:[pos="AA"]+ "man"
```

This will capture the adjectives found for each match in a captured group named "Adjectives".

BlackLab also supports numbered groups:

```
"an?|the" 1:[pos="AA"]+ "man"
```

Global constraints

If you tag certain tokens with labels, you can also apply "global constraints" on these tokens. This is a way of relating different tokens to one another, for example requiring that they correspond to the same word:

```
A:[] "by" B:[] :: A.word = B.word
```

This would match "day by day", "step by step", etc.