
textract Documentation

Release 0.5.1

Dean Malmgren

August 09, 2014

1	Currently supporting	3
2	Related projects	5
2.1	Command line interface	5
2.2	Python package	6
2.3	Installation	8
2.4	Contributing	9
2.5	Change Log	10
3	Indices and tables	13
	Python Module Index	15

As undesirable as it might be, more often than not there is extremely useful information embedded in Word documents, PowerPoint presentations, PDFs, etc—so-called “dark data”—that would be valuable for further textual analysis and visualization. While *several packages* exist for extracting content from each of these formats on their own, this package provides a single interface for extracting content from any type of file, without any irrelevant markup.

This package provides two primary facilities for doing this, the *command line interface*

```
textract path/to/file.extension
```

or the *python package*

```
# some python file
import textract
text = textract.process("path/to/file.extension")
```

Currently supporting

- `.doc` via `antiword`
- `.docx` via `python-docx`
- `.eml` via `python builtins`
- `.gif` via `tesseract-ocr`
- `.jpg` and `.jpeg` via `tesseract-ocr`
- `.json` via `python builtins`
- `.html` via `beautifulsoup4`
- `.odt` via `python builtins`
- `.pptx` via `python-pptx`
- `.pdf` via `pdftotext` (default) or `pdfminer`
- `.png` via `tesseract-ocr`
- `.ps` via `ps2text`
- `.txt` via `python builtins`

Please recommend other file types by either mentioning them on the [issue tracker](#) or by [contributing](#)

Related projects

Of course, `textract` isn't the first project with the aim to provide a simple interface for extracting text from any document. But this is, to the best of my knowledge, the only project that is written in python (a language commonly chosen by the natural language processing community) and is method agnostic about how content is extracted (more on this [here](#)). Here is a small sample of similar projects (feel free to add to the list):

- [Apache Tika](#) has [very similar, if not identical, aims as `textract`](#). It has impressive coverage of a wide range of file formats and is written in java.
- [`textract` \(node.js\)](#) has similar aims as this `textract` package (including an identical name! great minds...). It is written in node.js.
- [`pandoc`](#) is intended to be a document conversion (a much more difficult task!), but it does have [the ability to convert to plain text](#). It is written in Haskell.

Contents:

2.1 Command line interface

2.1.1 `textract`

Command line tool for extracting text from any document.

```
usage: textract [-h] [-o OUTPUT] [-m METHOD] [-v] filename
```

Positional arguments:

filename	Filename to extract text.
-----------------	---------------------------

Options:

-o=, --output=	output raw text in this file
-m=, --method=	specify a method of extraction for formats that support it
-v, --version	show program's version number and exit

Note: To make the command line interface as usable as possible, autocompletion of available options with `textract` is enabled by @kislyuk's amazing [argcomplete](#) package. Follow instructions to [enable global autocomplete](#) and you should be all set. As an example, this is also configured in the [virtual machine provisioning for this project](#).

2.2 Python package

This package is organized to make it as easy as possible to add new extensions and support the continued growth and coverage of textract. For almost all applications, you will just have to do something like this:

```
import textract
text = textract.process('path/to/file.extension')
```

to obtain text from a document.

For completeness, we also include here the documentation for specific file extension parsers as well as a few other essential bits in the `textract.exceptions` and `textract.shell` module.

2.2.1 `textract.parsers.doc_parser` module

```
textract.parsers.doc_parser.extract (filename, **kwargs)
```

Extract text from doc files using antiword.

2.2.2 `textract.parsers.docx_parser` module

```
textract.parsers.docx_parser.extract (filename, **kwargs)
```

Extract text from docx file using python-docx.

2.2.3 `textract.parsers.eml_parser` module

```
textract.parsers.eml_parser.extract (filename, **kwargs)
```

Extract text from email messages in .eml format. This gets the subject and all text from the contents.

2.2.4 `textract.parsers.gif_parser` module

2.2.5 `textract.parsers.html_parser` module

```
textract.parsers.html_parser.extract (filename, **kwargs)
```

Extract text from html file using BeautifulSoup4. Filter text to only show the visible parts of the page. Inspiration from [here](#).

2.2.6 `textract.parsers.jpg_parser` module

2.2.7 `textract.parsers.json_parser` module

```
textract.parsers.json_parser.extract (filename, **kwargs)
```

Extract all of the string values of a json file (no keys as those are, in some sense, markup). This is useful for parsing content from mongodb dumps, for example.

```
textract.parsers.json_parser.get_text (deserialized_json)
```

Recursively get text from subcomponents of a deserialized json. To enforce the same order on the documents, make sure to read keys of `deserialized_json` in a consistent (alphabetical) order.

2.2.8 textract.parsers.odt_parser module

`class textract.parsers.odt_parser.OpenDocumentTextFile (filepath)`

`textToString (element)`

`toString ()`

Converts the document to a string.

`textract.parsers.odt_parser.extract (filename, **kwargs)`

Extract text from open document files.

2.2.9 textract.parsers.pdf_parser module

`textract.parsers.pdf_parser.extract (filename, method=' ', **kwargs)`

Extract text from pdf files using method.

`textract.parsers.pdf_parser.extract_pdfminer (filename)`

Extract text from pdfs using pdfminer.

`textract.parsers.pdf_parser.extract_pdftotext (filename)`

Extract text from pdfs using the pdftotext command line utility.

2.2.10 textract.parsers.png_parser module

2.2.11 textract.parsers.pptx_parser module

`textract.parsers.pptx_parser.extract (filename, **kwargs)`

Extract text from pptx file using python-pptx

2.2.12 textract.parsers.ps_parser module

`textract.parsers.ps_parser.extract (filename, **kwargs)`

Extract text from postscript files using pstotext command.

2.2.13 textract.parsers.tesseract module

`textract.parsers.tesseract.extract (filename, **kwargs)`

Extract text from various image file formats using tesseract-ocr

2.2.14 textract.parsers.txt_parser module

`textract.parsers.txt_parser.extract (filename, **kwargs)`

Extract text from a .txt file

2.2.15 textract.cli module

`textract.cli.get_parser ()`

Initialize the parser for the command line interface and bind the autocompletion functionality

2.2.16 textract.exceptions module

exception `textract.exceptions.CommandLineError`

Bases: `exceptions.Exception`

The traceback of all `CommandLineError`'s is suppressed when the errors occur on the command line to provide a useful command line interface.

render (*msg*)

exception `textract.exceptions.ExtensionNotSupported` (*ext*)

Bases: `textract.exceptions.CommandLineError`

This error is raised with unsupported extensions

exception `textract.exceptions.MissingFileError` (*filename*)

Bases: `textract.exceptions.CommandLineError`

This error is raised when the file can not be located at the specified path.

exception `textract.exceptions.ShellError` (*command*, *exit_code*)

Bases: `textract.exceptions.CommandLineError`

This error is raised when a `shell.run` returns a non-zero exit code (meaning the command failed).

failed_message ()

is_uninstalled ()

uninstalled_message ()

exception `textract.exceptions.UnknownMethod` (*method*)

Bases: `textract.exceptions.CommandLineError`

This error is raised when the specified `--method` on the command line is unknown.

2.2.17 textract.shell module

`textract.shell.run` (*command*)

Run the specified shell command using Fabric-like behavior.

2.3 Installation

One of the main goals of `textract` is to make it as easy as possible to start using `textract` (meaning that installation should be as quick and painless as possible). This package is built on top of several python packages and other source libraries. Assuming you are using `pip` or `easy_install` to install `textract`, the [python packages](#) are all installed by default with `textract`. The source libraries are a separate matter though and largely depend on your operating system.

2.3.1 Ubuntu / Debian

There are two steps required to run this package on Ubuntu/Debian. First you must install some system packages using the [apt-get](#) package manager before installing `textract` from `pypi`.

```
apt-get install python-dev libxml2-dev libxslt1-dev antiword poppler-utils pstotext tesseract-ocr
pip install textract
```

2.3.2 OSX

These steps rely on you having [homebrew](#) installed as well as the [cask](#) plugin (`brew install caskroom/cask/brew-cask`). The basic idea is to first install [XQuartz](#) before installing a bunch of system packages before installing `textextract` from pypi.

```
brew cask install xquartz
brew install poppler antiword tesseract
pip install textextract
```

Note: [pstotext](#) is not currently a part of homebrew so `.ps` extraction must be enabled by manually installing from source.

Note: Depending on how you have python configured on your system with homebrew, you may also need to install the python development header files for `textextract` to properly install.

2.3.3 Don't see your operating system installation instructions here?

My appologies! Installing system packages is a bit of a drag and its hard to anticipate all of the different environments that need to be accomodated (wouldn't it be awesome if there were a system-agnostic package manager or, better yet, if python could install these system dependencies for you?!?!). If you're operating system doesn't have documentation about how to install the `textextract` dependencies, please [contribute a pull request](#) with:

1. A new section in here with the appropriate details about how to install things. In particular, please give instructions for how to install the following libraries before running `pip install textextract`:
 - [libxml2 2.6.21 or later](#) is required by the `.docx` parser which uses [lxml](#) via `python-docx`.
 - [libxslt 1.1.15 or later](#) is required by the `.docx` parser which users [lxml](#) via `python-docx`.
 - python header files are required for building `lxml`.
 - [antiword](#) is required by the `.doc` parser.
 - [pdftotext](#) is *optionally* required by the `.pdf` parser (there is a pure python fallback that works if `pdftotext` isn't installed).
 - [pstotext](#) is required by the `.ps` parser.
2. Add a requirements file to the [requirements directory](#) of the project with the lower-cased name of your operating system (e.g. `requirements/windows`) so we can try to keep these things up to date in the future.

2.4 Contributing

The overarching goal of this project is to make it as easy as possible to extract raw text from any document for the purposes of most natural language processing tasks. In practice, this means that this project should preferentially provide tools that correctly produce output that has words in the correct order but that whitespace between words, formatting, etc is totally irrelevant.

Importantly, this project is committed to being as agnostic about how the content is extracted as it is about the means in which the text is analyzed downstream. This means that `textextract` should support multiple modes of extracting text from any document and provide reasonably good defaults (defaulting to tools that tend to produce the correct word sequence).

Another important aspect of this project is that we want to have extremely good documentation. If you notice a type-o, error, confusing statement etc, please fix it!

2.4.1 Quick start

1. Fork and clone the project:

```
git clone https://github.com/YOUR-USERNAME/textract.git
```

2. Install [Vagrant](#) and [Virtualbox](#) and launch the development virtual machine:

```
vagrant plugin install iniparse
vagrant up && vagrant provision
```

On vagrant sshing to the virtual machine, note that the `PYTHONPATH` and `PATH` [environment variables](#) [have been altered in this virtual machine](#) so that any changes you make to `textract` in development are automatically incorporated into the command.

3. On the virtual machine, make sure everything is working by running the suite of functional tests:

```
./tests/run_functional_tests.sh
```

These functional tests are designed to be run on an Ubuntu 12.04 LTS server, just like the virtual machine and the server that runs the `travis-ci` test suite. There are some other tests that have been added along the way in the [Travis configuration](#). For your convenience, you can run all of these tests with:

```
./tests/run.py
```

Current build status:

4. Contribute! There are several [open issues](#) that provide good places to dig in. Check out the [contribution guidelines](#) and send pull requests; your help is greatly appreciated!

2.4.2 Style guidelines

As a general rule of thumb, the goal of this package is to be as readable as possible to make it easy for novices and experts alike to contribute to the source code in meaningful ways. Pull requests that favor cleverness or optimization over readability are less likely to be incorporated.

To make this notion of “readability” more concrete, here are a few stylistic guidelines that we recommend:

- write functions and methods that can [fit on a screen or two of a standard terminal](#) — no more than approximately 40 lines.
- unless it makes code less readable, adhere to [PEP 8](#) style recommendations — use an appropriate amount of whitespace.
- [code comments should be about *what* is being done, not *how* it is being done](#) — that should be self-evident from the code itself.

2.5 Change Log

This project uses [semantic versioning](#) to track version numbers, where backwards incompatible changes (highlighted in **bold**) bump the major version of the package.

2.5.1 latest changes in development

[will add changes here as they are made]

2.5.2 0.5.1

- several bug fixes, including:
 - documentation fixes
 - shell commands hanging on large files ([‘#33’](#))

2.5.3 0.5.0

- support for `.json` files ([#13](#) via [@anthonygarvan](#))
- support for `.odt` files ([#29](#) via [@christomitov](#))
- support for `.ps` files ([#25](#))
- support for `.gif`, `.jpg`, `.jpeg`, and `.png` files ([#30](#) via [@christomitov](#))
- several bug fixes, including:
 - improved fallback handling in `.pdf` parser if the `pdftotext` command line utility isn’t installed ([#26](#))
 - improved documentation for installation instructions on non-Ubuntu operating systems ([#21](#), [#26](#))
- several internal improvements, including:
 - cleaned up implementation of extension parsers to avoid magic

2.5.4 0.4.0

- support for `.html` files ([#7](#))
- support for `.eml` files ([#4](#))
- automated the documentation for the python package using sphinx-apidoc in docs/Makefile ([#9](#))

2.5.5 0.3.0

- support for `.txt` files, haha ([#8](#))
- fixed installation bug with not properly including requirements files in the manifest

2.5.6 0.2.0

- support for `.doc` files ([#2](#))
- support for `.pdf` files ([#3](#))
- several bug fixes, including:
 - fixing tab complete bug no file paths ([#6](#))
 - fixing tests to make sure the work properly on travis-ci

2.5.7 0.1.0

- Initial release, support for `.docx` and `.pptx`

Indices and tables

- *genindex*
- *modindex*
- *search*

t

- `textract.cli`, 7
- `textract.exceptions`, 8
- `textract.parsers.doc_parser`, 6
- `textract.parsers.docx_parser`, 6
- `textract.parsers.eml_parser`, 6
- `textract.parsers.gif_parser`, 6
- `textract.parsers.html_parser`, 6
- `textract.parsers.jpg_parser`, 6
- `textract.parsers.json_parser`, 6
- `textract.parsers.odt_parser`, 7
- `textract.parsers.pdf_parser`, 7
- `textract.parsers.png_parser`, 7
- `textract.parsers.pptx_parser`, 7
- `textract.parsers.ps_parser`, 7
- `textract.parsers.tesseract`, 7
- `textract.parsers.txt_parser`, 7
- `textract.shell`, 8