

# Outilex platform Graphical Interface - user guide

Olivier Blanc and Matthieu Constant

Oct. 11, 2006



# Contents

|          |                                              |           |
|----------|----------------------------------------------|-----------|
| <b>1</b> | <b>Introduction</b>                          | <b>1</b>  |
| <b>2</b> | <b>Getting started</b>                       | <b>3</b>  |
| 2.1      | System Requirements . . . . .                | 3         |
| 2.2      | Outilex directory . . . . .                  | 3         |
| 2.3      | Installation procedure . . . . .             | 4         |
| 2.4      | Launch Outilex . . . . .                     | 4         |
| 2.4.1    | Starting command . . . . .                   | 4         |
| 2.4.2    | UI general description . . . . .             | 5         |
| 2.4.3    | Projects . . . . .                           | 5         |
| <b>3</b> | <b>Dictionaries</b>                          | <b>7</b>  |
| 3.1      | DELA format . . . . .                        | 7         |
| 3.2      | XML format . . . . .                         | 8         |
| 3.2.1    | Description . . . . .                        | 8         |
| 3.2.2    | Lingdef . . . . .                            | 9         |
| 3.3      | Operations . . . . .                         | 10        |
| 3.3.1    | Editing a dictionary . . . . .               | 10        |
| 3.3.2    | Converting DELA in XML . . . . .             | 10        |
| 3.3.3    | Indexing dictionaries . . . . .              | 11        |
| 3.3.4    | Add selected dictionary to project . . . . . | 12        |
| <b>4</b> | <b>Grammars</b>                              | <b>13</b> |
| 4.1      | A simple graph . . . . .                     | 13        |
| 4.2      | Subgraphs . . . . .                          | 14        |
| 4.3      | Ouputs and weights . . . . .                 | 14        |
| 4.4      | Normalization graphs . . . . .               | 16        |
| 4.5      | Decoration graphs . . . . .                  | 17        |
| <b>5</b> | <b>Text FSA</b>                              | <b>19</b> |

|          |                                    |           |
|----------|------------------------------------|-----------|
| <b>6</b> | <b>Text processing</b>             | <b>23</b> |
| 6.1      | Text segmentation . . . . .        | 23        |
| 6.2      | Dictionary application . . . . .   | 24        |
| 6.3      | Normalize text automaton . . . . . | 24        |
| 6.4      | Grammar application . . . . .      | 24        |
| 6.5      | Locate pattern . . . . .           | 24        |
| <b>7</b> | <b>C++ Programs</b>                | <b>25</b> |
| 7.1      | apply-dic . . . . .                | 25        |
| 7.2      | concordancer . . . . .             | 25        |
| 7.3      | concordancer2 . . . . .            | 26        |
| 7.4      | decure-fsa . . . . .               | 27        |
| 7.5      | dela-index . . . . .               | 27        |
| 7.6      | dic-lookup . . . . .               | 27        |
| 7.7      | delaf2xml.sh . . . . .             | 27        |
| 7.8      | dic-index . . . . .                | 28        |
| 7.9      | make-concord-html . . . . .        | 28        |
| 7.10     | make-wrtn . . . . .                | 28        |
| 7.11     | tfsa-copy . . . . .                | 28        |
| 7.12     | tfsa2dot . . . . .                 | 29        |
| 7.13     | tokenization . . . . .             | 29        |
| 7.14     | transduct-fsa . . . . .            | 29        |
| 7.15     | wrtn-flatten . . . . .             | 29        |
| 7.16     | u82i . . . . .                     | 30        |
| 7.17     | u82u16 . . . . .                   | 30        |
| 7.18     | u1628 . . . . .                    | 30        |
| 7.19     | wrtn-txt-transduct . . . . .       | 30        |
| <b>8</b> | <b>Appendix: a text automaton</b>  | <b>31</b> |

# Chapter 1

## Introduction

Outilex is a 4-year-research project funded by the French Industry Ministry. It gathers 4 academic institutions and 6 industrial organizations:

- Centre de l'Energie Atomique (CEA)
- Institut Gaspard Monge (IGM), Universite de Marne-la-Vallee (coordinator)
- Laboratoire de Psychologie et Neurosciences de la Cognition de l'Université de Rouen
- Laboratoire d'Informatique de Paris 6 (LIP6)
- Langage Communication Information (LCI)
- Lingway
- LORIA
- Systran
- Thales Com
- Thales R&D

Started in October 2002, this project aims at developing a platform devoted to Natural Language Processing.



## Chapter 2

# Getting started

### 2.1 System Requirements

To compile the C++ programs, you need

- compiler `gcc/g++` with version  $\geq 4$
- tool `jam` (from Perforce)
- library `libxml` (standard XML library)
- library `ICU` from IBM
- library `C++ Boost` ([www.boost.org](http://www.boost.org)), compiled with `ICU`

To compile the Systran C++ tokenization module, you need an old version of `flex` 2.5.4 (package `flex-old`). Newer versions do not work.

To run the User Interface, you need the `Java Run time Environment` (1.5).

### 2.2 Outilex directory

Outilex directory includes the following files and directories:

- file `README.txt` (to get started)
- file `install-outilex` (script to compile and install C++ programs)
- file `outilexUI.jar` (to run user interface)
- file `clean-outilex` (script to clean compiled programs)

- directory **bin** (C++ compiled programs)
- directory **data** (some linguistic data provided with the platform)
- directory **docs** (documentation)
- directory **lingdef** (linguistic definitions of the set of tags used in dictionaries and graphs)

This directory also contains a log file (**outilex.log**) that includes all commands that have been launched from the interface. This can help users getting used with the syntax of the commands of the different programs.

## 2.3 Installation procedure

To install Outilex platform, you need to follow the steps described below:

- Go to Outilex directory;
- Run compilation by typing:

```
./install-outilex
```

**Warning:** You might have to change some environment parameters depending on you local system.

- Set LINGDEF environment variable by typing:

```
LINGDEF=${OUTILEX_HOME}/lingdef/french/lingdef.xml  
export LINGDEF
```

with `${OUTILEX_HOME}` the path of Outilex directory

To avoid doing this operation everytime you want to run Outilex-platform, you should put these commands in your `.bashrc` file.

**Note:** this operation is temporary and should be removed in the next versions of the Outilex platform.

## 2.4 Launch Outilex

### 2.4.1 Starting command

Outilex platform User Interface (UI) can be launched by typing:

```
java -jar outilexUI.jar
```

### 2.4.2 UI general description

The UI is composed of:

- a menu (on top), (Tool bar, coming soon...)
- a process/resource panel (on the left): to create personal processing chain with available linguistic resources,
- a content panel (on the right): to display linguistic resources and processing results

**Important note:**

Many functionalities can be run via popup menus (right-click on the mouse). Double-clicking on a resource (in the left panel), makes it display on the rightpanel.

### 2.4.3 Projects

Outilex platform works with a system of project. Each project is composed of a set of resources (texts, dictionaries and grammars). The Menu **Project** allows users to create, open and save projects. A project is associated with one language. This language selects a linguistic definition file. For example, if 'french' is the project language, the set of linguistic tags that will be used in the processings is defined in the file `lingdef/french/lingdef.xml`.



## Chapter 3

# Dictionaries

Dictionaries are sets of lexical entries associated with morphological, syntactic and semantic information. Lexical entries are either simple words or multiword units. Outilex platform allows users to edit their own dictionaries. There are several formats:

- an editable textual format: DELA format encoded in UTF-8 (extension `.dic`)
- an exchange format in XML (extension `.dic.xml.gz`)
- a binary format used by programs (extension `.idx`)

The different operations on dictionaries are gathered in the **Dictionary** menu of the platform.

### 3.1 DELA format

The DELA format has been defined in [?, ?].

#### Syntax of an entry

An entry is defined on a single line as it is shown in the following examples:

```
car,.N+Conc:s/this is an example
eats,eat.V:P3s
Tony Blair,.N+Npr+Hum:ms
sincerely,.ADV
```

- The first element is the inflected form and is **obligatory** (**car** and **eats**).
- The second element (between symbols ',' and '.') is the lemma and is optional (e.g. **eat**). If it is not present, the lemma is considered to be the same as the inflected form (e.g. **car**).
- The third element is the part-of-speech and is **obligatory** (e.g. **ADV** for adverb, **N** for noun).
- The elements following symbol '+' are syntactic and semantic information and are optional (e.g. **Conc** for concrete, **Npr** for proper name, **Hum** for human).
- The character sequence following symbol ':' is a set of morphological information and are optional; each character stands for a piece of information (e.g. **P3s** stands for present [P] at the third person [3] of singular [s]).
- The sequence following symbol '/' is an optional comment (e.g. **this is an example**).

The tagset is free, as long as the writer follows the syntax defined above. This editable dictionaries are encoded in UTF-8 and their file extension is **.dic**.

## 3.2 XML format

### 3.2.1 Description

Outilex uses an UTF-8 XML exchange format (file extension **.dic.xml**). Tagset is defined in the **lingdef** (cf. section 3.2.2). All tags used must be defined in the **lindef** file.

Hereby is an example of an entry :

```
<entry>
<lemma>abaissable</lemma>
<pos name='adj' />
<inflected>
  <form>abaissable</form>
  <feat name='gender' value='masculine' />
  <feat name='number' value='singular' />
</inflected>
</entry>
```

```

</inflected>
<inflected>
  <form>abaissable</form>
  <feat name='gender' value='feminine' />
  <feat name='number' value='singular' />
</inflected>
<inflected>
  <form>abaissables</form>
  <feat name='gender' value='masculine' />
  <feat name='number' value='plural' />
</inflected>
<inflected>
  <form>abaissables</form>
  <feat name='gender' value='feminine' />
  <feat name='number' value='plural' />
</inflected>
</entry>

```

To avoid memory space feeding, XML dictionaries are compressed and their file extension is `.dic.xml.gz`.

### 3.2.2 Lingdef

The Lingdef file defines the tagset that can be used in Outilex XML dictionaries. It is also encoded in XML. This file is located in the directory `<language>` (e.g. `french`) of the directory `lingdef`.

Hereby is an example of the definition of the tagset used for nouns:

```

<!-- nouns -->

<attrtype name='nounsubcat' type='enum'>
  <value name='pred' alias='Pred' />
  <value name='conc' alias='Conc,concret' />
  <value name='abst' alias='Abst,abstract,abs' />
  <value name='hum' alias='Hum,human' />
  <value name='anl' alias='Anl,animal' />
  <value name='tps' alias='Tps,temporal' />
  <value name='top' alias='Top,toponym' />
  <value name='unit' alias='Unit' />
  <value name='num' alias='Nnum,numeral' />
  <value name='dnom' alias='Dnom,detnom' />

```

```

</attrtype>

<attrtype name='collective' type='bool'>
  <true alias='Coll' />
</attrtype>

<attrtype name='proper' type='bool'>
  <true alias='pr,Pr,Prenom' />
</attrtype>

<pos name='noun' cutename='N'>
  <attribute name='subcat' type='nounsubcat' shortcut='yes' />
  <attribute name='gender' type='gender' shortcut='yes' />
  <attribute name='number' type='number' shortcut='yes' />
  <attribute name='proper' type='proper' default='false' shortcut='yes' />
  <attribute name='coll' type='collective' shortcut='yes' />
</pos>

```

### 3.3 Operations

#### 3.3.1 Editing a dictionary

For editing a new or existing dictionary, you need to click on items **new** or **open** in the menu **Dictionary**. Then a text editor containing the dictionary will appear in the content part of the UI. After having edited the dictionary, you can save it by clicking on **save** or **save as** in the same menu. Dictionaries are saved in UTF-8.

#### 3.3.2 Converting DELA in XML

Converting a DELA dictionary into an XML one requires the definition of the file **delaf-corresp** that defines the correpondance between tags used in the DELA dictionary and tags used in the XML dictionary. This file is located in the directory **<language>** (e.g. **french**) of the directory **lingdef**.

Below is a sample of this file:

|         |        |
|---------|--------|
| POS A   | adj    |
| POS ADV | adverb |
| POS DET | det    |
| POS N   | noun   |

```

...

flex f PREPDET,DET,PRO,A,N,V form gender e feminine
flex m PREPDET,DET,PRO,A,N,V form gender e masculine
flex p PREPDET,DET,PRO,A,N,V form number e plural
flex s PREPDET,DET,PRO,A,N,V form number e singular

...

synt d A lemma postpos b true
synt g A lemma antepos b true
....

```

Each line defines a correspondance for a given code used in the DELA. It contains several fields. The first field define the type of code that is described (POS for part-of-speech; **flex** for inflection code ; **synt** for syntactical or semantic code). The line

```
POS N      noun
```

indicates that the part-of-speech N in the DELA dictionary (meaning noun) will be used as **noun** in the XML one.

The line

```
flex s PREPDET,DET,PRO,A,N,V form number e singular
```

means that **s** is an infection code that is only valid for part-of-speech PREPDET, DET, PRO, A, N and V. The sequence **form number e singular** shows that it is defined as an attribute of the inflected form (**form**), where **number** is the name of the attribute, **e** is the type of the attribute (**e** for enumeration) and **singular** is the value of the attribute.

For converting a DELA dictionary into an XML dictionary, you need to click on item **Transcode/DELA -> XML** in the menu **Dictionary**. The selected DELA dictionary will then be converted in a compressed XML format and be displayed in the content part of the UI. (Program used: **delaf2xml.sh**, cf. section 7.7).

### 3.3.3 Indexing dictionaries

To be used in the Outilex text processes, the dictionaries must be indexed into binary files that represent the dictionaries in the form of minimized finite-state transducers. These files have the extension **.idx**.

To do so, you need to click on item `indexing` in the menu `Dictionary`. The C++ program that is launched is either `dic-index` for compressed XML dictionaries (cf. section 7.8) or `dela-index` for DELA dictionaries (cf. section 7.5).

This process through the interface will produce a file `<name>.idx` if input dictionaries are named `<name>.dic` or `<name>.dic.xml.gz`.

### 3.3.4 Add selected dictionary to project

You can add the selected dictionary to your current project by clicking on item `Add to project` in the menu `Dictionary`, ONLY IF IT HAS ALREADY BEEN INDEXED.

## Chapter 4

# Grammars

Grammars used in Outilex are equivalent to Recursive Transition Networks [?]. They are in the form of graphs. Their symbols used can be lexical values, lexical masks or call to sub-graphs. They can also contain outputs and weights. Outilex platform includes a graph editor developed from Unitex sources [?]. This graph editor generates graph encoded in XML, with the extension `.xgrf`. Some example graphs are provided with the platform.

**Temporary, for more details on graph edition, see the Unitex manual (file `manuelunitex.pdf`)**

### 4.1 A simple graph

Outilex graphs are oriented graphs represented with oriented "rectangular" boxes filled with text, and transitions that relate these boxes. There exist two obligatory special boxes: an initial one (here, in the form of an arrow) and a final one (in the form of a circle). Generally, they are read from left to right, from the initial box to the final box. An example is given in figure 4.1. The text of a box is set in an `input` text field. For instance, the first box on the left of the example graph is defined with the following text:

```
lundi+mardi+mercredi+jeudi+vendredi+samedi+dimanche
```

where symbol '+' stands the union symbol of a regular expression and is represented in the box with a line breaking. This graph therefore recognizes a text sequence like `mardi prochain`.

For technical details on graph edition, cf. Unitex manual.

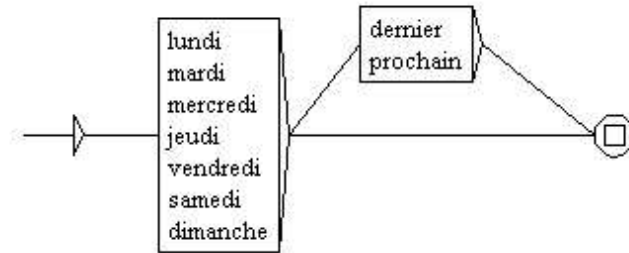


Figure 4.1: A simple graph

## 4.2 Subgraphs

A call to a subgraph can be designed in a graph box by typing the name of the subgraph (without extension) preceded by the symbol ':' in the **input** text field. The subgraph should be in the same directory<sup>1</sup>. The box calling a subgraph is greyed to distinguish them from "normal" boxes, as illustrated in figure 4.2. For instance, the box defined by the following **input** text

```
:subgraph
```

is a call to the graph `subgraph.xgrf`.



Figure 4.2: A call to a subgraph

## 4.3 Outputs and weights

If you want to associate an output with a box, you need to use the **output** text field of the box. As in the input, symbol '+' can play the role of a "line breaker". Weights can also be added with the following syntax in the **output** text field:

```
<output>/<weight>
```

---

<sup>1</sup>Future improvements should authorize call to subgraphs in other directories

where `<output>` is a string defining the output associated with an input, `<weight>` is a positive real number. For instance,

Noun/3.0

The weight of a path of a graph is the sum of all its weights. The weight is optional : by default, if the weight of an input label is missing, its value is 0. Then the **output** text field can simply be defined by the text:

`<output>`

For instance, the graph in figure 4.3 recognizes sequences **noun-adj**. But when a compound noun (**noun+comp**) is also recognized on the same sequence of text, a priority is given to this last analysis because it is assigned a weight of 1 (0 for the other).

As a consequence, the symbol `'/'` is a specialized symbol to delimit `<output>` and `<weight>`. So, when using such a symbol in the `<output>` string, field `<weight>` is obligatory and user should put an actual weight or an empty string following the last symbol `'/'`. For instance in figure 4.4, the output of the last box is defined with the following text :

`</date>/`

where symbol `'/'` is present at the end because field `<output>` contains such a symbol ; the associated weight is the default weight (0) because `<weight>` is the empty string.

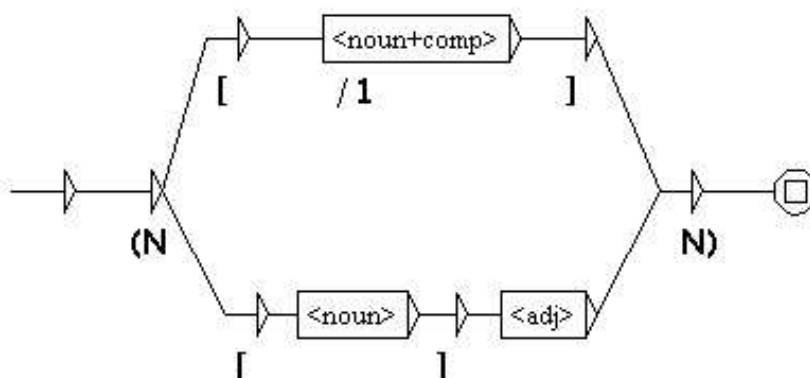


Figure 4.3: Use of weights

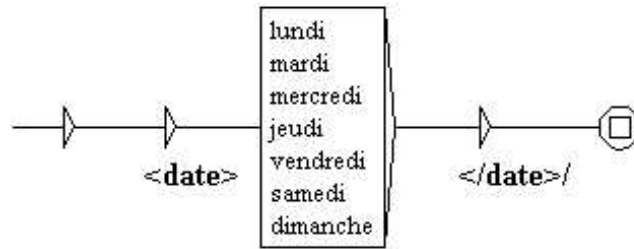


Figure 4.4: Special outputs

## 4.4 Normalization graphs

A normalization graph is a graph such that when applied to a text-automaton it normalizes some sequences like **de** as shown in the following example:

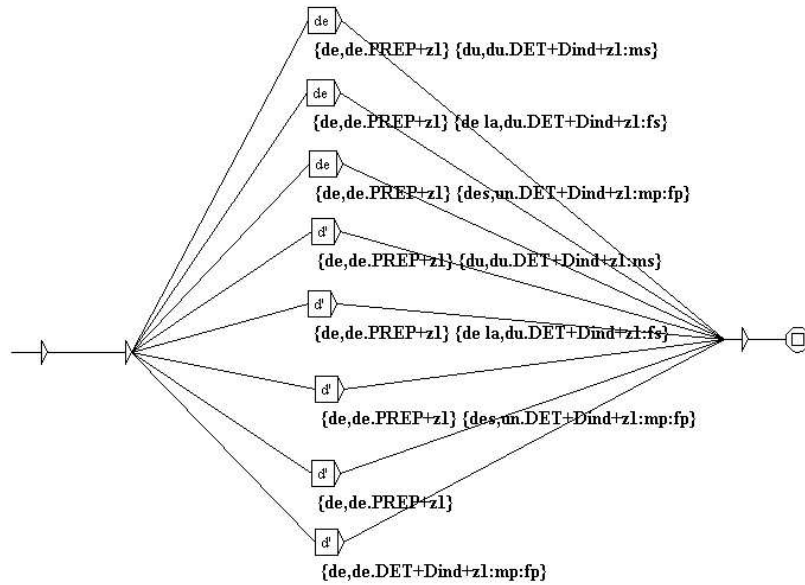


Figure 4.5: Normalization graph

## 4.5 Decoration graphs

A decoration graph is a graph such that when applied to a text automaton, new transitions corresponding to new analyses are added to the initial text automaton.

Below are two examples:

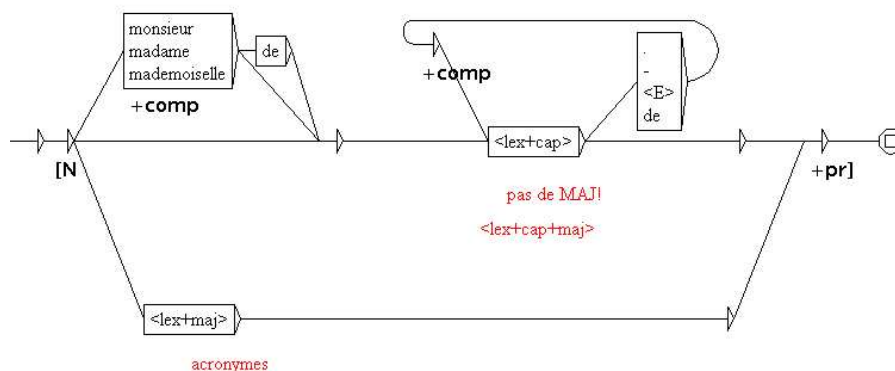


Figure 4.6: Recognition of named entities

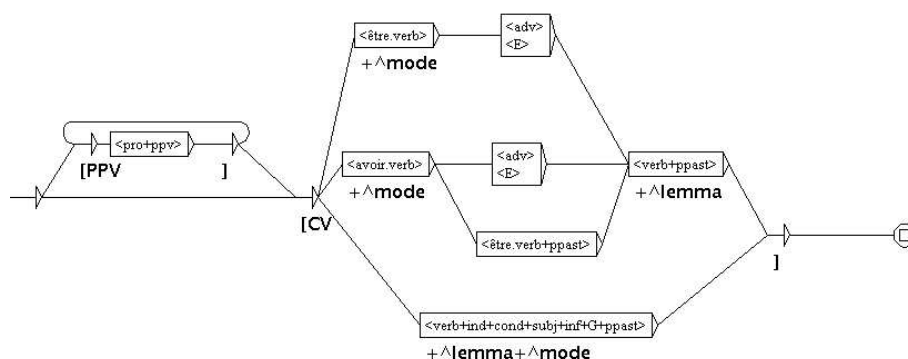


Figure 4.7: Recognition of verbal complexes

These graphs have a special format. Linguistic entities that have to be analysed must be delimited in the graph by square brackets in the output like in the graphs shown above. The part-of-speech (e.g. N for "noun" in figure 4.6, CV for "verbal complex" in figure 4.7 PPV for "preverbal pronoun" in figure 4.7) to be assigned to these entities must be put just after the

opening square brackets. Attributes to this part-of-speech can also be added by adding outputs with the following syntax `+<attribute>` (e.g. `+Npr` means "proper name" in figure 4.6). For instance, the graph in figure 4.6 recognizes named entities which are tagged `N+Npr`.

A complex entity can inherit from attributes from its elements (e.g. lemma, mode, gender, and so on.). This can be indicated in the decoration graph as an output with the following syntax `+^<attribute>`. Such an example is shown in figure 4.7: when recognized in a text, the pattern `<avoir.verb> <verb+ppast>`<sup>2</sup> is analysed as a `CV` (verbal complex) whose mode is the mode of `avoir` and whose lemma is the one of the verb at the past participle. For instance, the sequence `a lu` would be analysed as a `CV` whose mode is `ind` (for indicative) and whose lemma is `lire`.

---

<sup>2</sup>the verb `avoir` followed by a verb at the past participle

## Chapter 5

# Text FSA

The Text FSA is used to represent text ambiguity. For each sentence, there is an automaton that represents its possible analyses. Grammar application programs all use it as input. Given the text containing two sentences:

**Cinq marins sont toujours portés disparus, seul un membre de l'équipage a été secouru. Il a été entendu par les enquêteurs.**

After application of the French dictionary provided for the project by IGM, the text is represented as a set of two automata, each automaton associated with a sentence. A graphical example of the second sentence is given in figure 5.1.

By default, the resulting text automaton is a binary file. Nevertheless, Outilex provides a converter into a XML file. Extracts of such file is provided below (**tfssa-copy**). The complete file is given in appendice. The element **text-fsa** contains first an element **lexic** and a set of elements **sfsa**. **lexic** includes the set of lexical entries that will be used in the labels of the automata transitions. Lexical entries are represented the same way as defined in the **lingdef** file. Elements **sfsa** represent sentence automata, which are defined as a set of states (**q**) from which some transitions (**tr**) start.

```
<?xml version="1.0"?>
<text-fsa sz="2">
  <lexic sz="79">
    <lex>
      <form>Cinq</form>
      <lem>cinq</lem>
      <pos v="lex"/>
      <f n="case" v="cap"/>
```

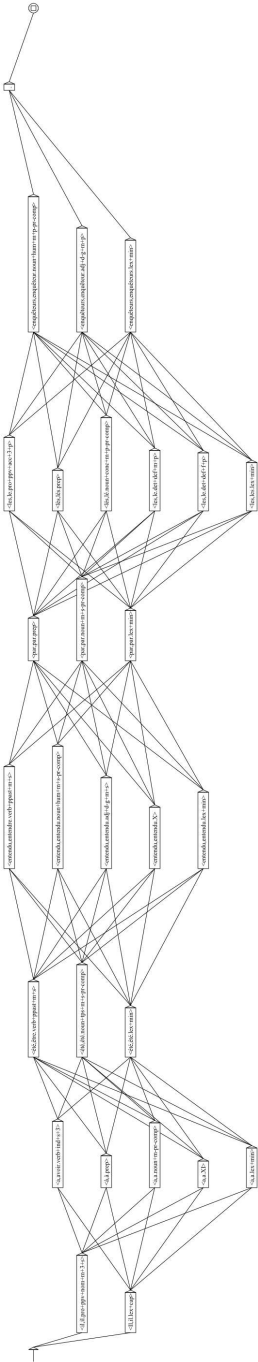


Figure 5.1: A sentence automaton

```

</lex>
<lex>
  <form>cinq</form>
  <lem>cinq</lem>
  <pos v="det"/>
  <f n="gender" v="f"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>cinq</form>
  <lem>cinq</lem>
  <pos v="det"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>cinq</form>
  <lem>cinq</lem>
  <pos v="noun"/>
  <f n="subcat" v="num"/>
  <f n="gender" v="m"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>marins</form>
  <lem>marins</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
  .....
</lexic>
<sfsa sz="19">
  <txt>Cinq marins sont toujours portés disparus,
    seul un membre de l'équipage a été secouru. </txt>
  <q id="0" pos="0">
    <tr lbl="0" to="1"/>
    <tr lbl="1" to="1"/>
    <tr lbl="2" to="1"/>
    <tr lbl="3" to="1"/>

```

```
</q>
<q id="1" pos="0">
  <tr lbl="4" to="2"/>
  <tr lbl="5" to="2"/>
  <tr lbl="6" to="2"/>
</q>
<q id="2" pos="0">
  <tr lbl="7" to="3"/>
  <tr lbl="8" to="3"/>
</q>
.....
<q id="18" pos="0" f="1"/>
</sfsa>
.....

</text-fsa>
```

## Chapter 6

# Text processing

Outilex platform allows users to process texts using linguistic resources such as dictionaries and grammars. The left part of the UI permits to create your own desired chain.

Processing a text is:

- segment text in tokens and sentences (check box **segmentation**),
- applying dictionaries on segmented text and obtaining a text automaton representing the possible analyses for each sentence (check box **apply dictionaries**),
- normalizing text-fsa by applying normalization graphs (check box **normalize**).
- apply a cascades of grammars in the form of graphs on the text automaton, resulting to a new text automaton (check box **apply graph cascades**),
- applying a grammar to the text automaton to obtain a concordance or a modified text (e.g. an annotated text) (check box **locate pattern**).

### 6.1 Text segmentation

The text segmentation process creates a directory associated to the text (`<text>.dir`) and outputs a segmented text `<text>.segmentation` put in this directory.

## 6.2 Dictionary application

You must insert and select dictionaries you need (click on button **more**, click on button **less** to close). The process will generate a text automaton, `<text>.fsa`, in the text directory. A copy of it is also made (file `<text>-0.fsa`)

## 6.3 Normalize text automaton

This operation must be run after dictionary application and text-automaton construction. It applies the graph `Norm.xgrf` in the directory `lingdef/<language>` where `<language>` is the current language. This process generates a new "normalized" text-automaton (file `<text>-norm.fsa`).

## 6.4 Grammar application

You need to define a list of graphs that will be applied in cascades on `<text>.fsa`. Each iteration `j` will generate a new text automaton `<text>-j.fsa`. The final automaton is `<text>-final.fsa`. A copy of it is made in file `<text>.fsa`;

`<text>.fsa` is actually the current text automaton to be processed.

## 6.5 Locate pattern

You need to select a graph to be applied and the type of result you want.

## Chapter 7

# C++ Programs

The Outilex platform is made of a set of independant C++ programs. This chapter defines their different prototypes.

### 7.1 apply-dic

```
apply-dic -dic <dic1> [<prio1>] [-dic <dic2> [<prio2>] ...]  
        [-imaj] [-icase] [-imark] [-l <lingdef>] [-o <out>] <tokfile>
```

This program applies a set of dictionaries <dicj> (extension .idx) with different priorities <prioj> (real numbers, by default, 10) to a segmented text <tokfile>. It outputs a text-fsa <out> (by default, the name of the text file with the extension .fsa). It uses a lingdef file <lingdef>. Options could be:

- -imaj: ignore case in texts but not in dictionaries;
- -icase: ignore case in dictionaries and in texts;
- -imark: ignore diacritics in texts and in dictionaries.

### 7.2 concordancer

```
concordancer -l <lingdef> -gram <gram> [-v] [-longest-match]  
        [-tags] [-tree] [-w] [-m] [-ipath] [-iout] [-o <outputres>] <txtfsa>
```

with options :

<txtfsa> : input text fsa

-gram <gram> : wrtn grammar to apply  
 -o <concord> : name of the resulting concordance index file (default to concord.idx)  
 -longest-match : keep only longest matching sequences  
 -tags : display morpho-syntactic tags  
 -tree : display syntactic tree  
 -w : display weights of matching sequences  
 -m : merge grammar's outputs into concordances  
 -all : shortcut for : -tags -tree -w -m  
 -ipath : keep only one concordance for the same text segment (can be a lot faster for ambiguous grammars)  
 -v : verbose mode (for debugging)  
 It applies a compiled wrtn grammar <gram> (extension .wrtn) to a text fsa <txtfsa> and saves the matching sequences index into a file <outputres> (default to concord.idx), which can be processed by make-concord-html. There exist different options that are described above.

### 7.3 concordancer2

concordancer -l <lingdef> -gram <gram> [-v] [-longest-match]  
 [-tags] [-tree] [-w] [-m] [-ipath] [-iout] [-o <outputres>] <txtfsa>

<txtfsa> : input text fsa  
 -l <lingdef> : lingdef to be used  
 -gram <gram> : wrtn grammar to apply  
 -o <concord> : name of the resulting concordance index file (default to concord.xml)  
 -longest-match : keep only longest matching sequences  
 -tags : display morpho-syntactic tags  
 -tree : display syntactic tree  
 -w : display weights of matching sequences  
 -m : merge grammar's outputs into concordances  
 -all : shortcut for : -tags -tree -w -m  
 -timeout <s> : specify a maximum amount of time (in seconds) to spend to parse a sentence  
 -ipath : keep only one concordance for the same text segment (can be a lot faster for ambiguous grammars)  
 -v : verbose mode (for debugging)

It applies a wrtn grammar to a text fsa and saves the matching sequences index into a file (default to concord.xml), which can be processed by make-concord-html.

## 7.4 decore-fsa

```
decore-fsa -l <lingdef> -rtn <fst> [-v] [-ipath] [-iout] [-o <outputres>]
<txtfsa>
```

This program applies a compiled decoration grammar **fst** (extension **.wrt**) to the text-fsa **<txtfsa>**. It outputs a new version of **txtfsa**, **<outputres>**, with new transitions whenever new analyses have been found by applying grammar **fst**. It requires the lingdef file **<lingdef>**.

## 7.5 dela-index

```
dela-index <dela> -corresp <corresp> [-r <ratio>] [-o <index>]
```

This program compresses an UTF-8 DELA dictionary **<dela>** into into an IDX dictionary using the tag correspondance file **<corresp>**. The output file is **<index>**.

## 7.6 dic-lookup

```
dic-lookup [-imaj] [-imark] [-icase] -dic <diconame> [ req1 [req2
...]]
```

This program searches the entries of the words **req1 req2 ...** in the dictionary **diconame**. Options could be:

- -imaj: ignore case in words but not in dictionaries;
- -icase: ignore case in dictionaries and in words;
- -imark: ignore diacritics in words and in dictionaries.

## 7.7 delaf2xml.sh

```
delaf2xml.sh -c <corresp> <dela>
```

This program converts an UTF-8 DELA dictionary `<del>` (extension `.dic`) into a compressed XML dictionary (`<del>.xml`) using the tag correspondance file `<corresp>`.

## 7.8 dic-index

```
dic-index [-validate] [-ratio <r>] <dicofile>
```

This program compresses the dictionary `<dicofile>` (extension `.dic.xml.gz`) into an IDX dictionary. The output file is the name of `<dicofile>` with the extension (`.idx`). For instance, if `<dicofile>` is `dico.dic.xml.gz`, the output would be `dico.dic.idx`.

## 7.9 make-concord-html

```
make-concord-html <concordidx>  
  [-left <left-size>] [-right <right-size>] [-o <res>] [-dontsort]
```

It constructs an html concordance from the index concordance file `<concordidx>` with a left context of `<left-size>` characters (default: 50) and a right context of `<right-size>` characters (default: 80). Optionally, the concordance can be put in the text order (option `-dontsort`); by default, it is sorted. The result is put in the file `<res>` (default: `concord.html`).

## 7.10 make-wrtn

```
make-wrtn [-l <lingdef>] <axiom>
```

This program compiles the grammar defined by the main XGRF graph `<axiom>` and its sub-graphs into a unique XML file representing a Weighted Recursive Transition Network (WRTN) with the extension `.wrtn`. It uses the lingdef file `<lingdef>` to interpret semantically the symbols of the graphs. For instance, if `<axiom>` is `main.xgrf`, the output would be `main.wrtn`.

## 7.11 tfsa-copy

```
tfsa-copy [-o <out>] [-gz] [-f <oformat>] <txtfisa>
```

This program converts a binary text fsa `<txtfsa>` into XML (if option `-f` is `xml`). Option `-gz` indicates that the result will be gzipped.

## 7.12 tfsa2dot

```
tfsa2dot -l <lingdef> [-o <output>] <txtfsa> -n <sentenceno>
```

This program converts the `<sentenceno>`-th sentence of text-fsa `txtfsa` (extension `.fsa`) into the file `<output>` describing an automaton with the DOT format, using the lingdef file `<lingdef>`. By default, the output file is named `sentence-<sentenceno>.dot`.

## 7.13 tokenization

```
tokenization <text>
```

This program is used to segment a text `<text>` in tokens, sentences and paragraphs. `<text>` is the input text file name and can be either in TXT format or HTML format. Outputs are `<text>.segmentation`, `<text>.tokenization` and `<text>.postfilter`.

## 7.14 transduct-fsa

```
transduct-fsa -l <lingdef> -gram <fst>  
[-lgst] [-ipath] [-iout] [-dontsurf] [-o <outputres>] <txtfsa>
```

This program applies a normalization wrtn transducer `<fst>` on a text-fsa `<txtfsa>`. It produces a new text-fsa `<outputres>`.

## 7.15 wrtn-flatten

```
wrtn-flatten <rtn> [-maxdepth <N>]
```

This program flattens a compiled grammar `<rtn>` into a finite state automaton (or transducer). Whenever not possible, it makes an approximation by limiting the maximum depth (`<N>`).

**7.16 u82i****u82i**

This program converts the utf-8 encoded input into ASCII output.

**7.17 u82u16****u82u16**

This program converts the utf-8 encoded input into a little endian unicode output.

**7.18 u1628****u1628**

This program converts the input encoded in little-endian unicode encoded into an output encoded in utf-8.

**7.19 wrtn-txt-transduct**

```
wrtn-txt-transduct -l <lingdef> -gram <fst>
[-ipath|-iout] [-html|-txt] [-m|-r|-i] [-o <outputres>] <txtfsa>
```

This program applies wrtn transducer **<fst>** to a text fsa **<txtfsa>** and generates a new text from the original one depending on the chosen option:

- -m for merging outputs in the text when finding matching sequences;
- -r for replacing matching sequences by the associated outputs in the new text;
- -i for ignoring outputs: new text is the original one (USELESS!!!).

The output text **<outputres>** can be either in HTML (**-html**) or in TXT (**-txt**). It requires the lingdef file **<lingdef>**.

## Chapter 8

# Appendice: a text automaton

```
<?xml version="1.0"?>
<text-fsa sz="2">
  <lexic sz="79">
    <lex>
      <form>Cinq</form>
      <lem>cinq</lem>
      <pos v="lex"/>
      <f n="case" v="cap"/>
    </lex>
    <lex>
      <form>cinq</form>
      <lem>cinq</lem>
      <pos v="det"/>
      <f n="gender" v="f"/>
      <f n="number" v="p"/>
    </lex>
    <lex>
      <form>cinq</form>
      <lem>cinq</lem>
      <pos v="det"/>
      <f n="gender" v="m"/>
      <f n="number" v="p"/>
    </lex>
  </lex>
</text-fsa>
```

```

<form>cinq</form>
<lem>cinq</lem>
<pos v="noun"/>
<f n="subcat" v="num"/>
<f n="gender" v="m"/>
<f n="proper" v="false"/>
<f n="compound" v="false"/>
</lex>
<lex>
  <form>marins</form>
  <lem>marins</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>
  <form>marins</form>
  <lem>marin</lem>
  <pos v="adj"/>
  <f n="postpos" v="true"/>
  <f n="antepos" v="false"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>marins</form>
  <lem>marin</lem>
  <pos v="noun"/>
  <f n="subcat" v="hum"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>sont</form>
  <lem>sont</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>

```

```

<form>sont</form>
<lem>être</lem>
<pos v="verb"/>
<f n="mode" v="ind"/>
<f n="number" v="p"/>
<f n="pers" v="3"/>
</lex>
<lex>
  <form>toujours</form>
  <lem>toujours</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>
  <form>toujours</form>
  <lem>toujours</lem>
  <pos v="adv"/>
</lex>
<lex>
  <form>portés</form>
  <lem>portés</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>
  <form>portés</form>
  <lem>porté</lem>
  <pos v="adj"/>
  <f n="postpos" v="true"/>
  <f n="antepos" v="false"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>portés</form>
  <lem>porté</lem>
  <pos v="noun"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
  <f n="proper" v="false"/>

```

```

    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>portés</form>
    <lem>porter</lem>
    <pos v="verb"/>
    <f n="mode" v="ppast"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
  </lex>
  <lex>
    <form>disparus</form>
    <lem>disparus</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>disparus</form>
    <lem>disparu</lem>
    <pos v="adj"/>
    <f n="postpos" v="true"/>
    <f n="antepos" v="false"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
  </lex>
  <lex>
    <form>disparus</form>
    <lem>disparu</lem>
    <pos v="noun"/>
    <f n="subcat" v="hum"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>disparus</form>
    <lem>disparaître</lem>
    <pos v="verb"/>
    <f n="mode" v="ppast"/>

```

```

    <f n="gender" v="m"/>
    <f n="number" v="p"/>
  </lex>
  <lex>
    <form>disparus</form>
    <lem>disparaître</lem>
    <pos v="verb"/>
    <f n="mode" v="ind"/>
    <f n="number" v="s"/>
    <f n="pers" v="1|2"/>
  </lex>
  <lex>
    <form>,</form>
    <lem>,</lem>
    <pos v="punc"/>
  </lex>
  <lex>
    <form>seul</form>
    <lem>seul</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>seul</form>
    <lem>seul</lem>
    <pos v="adj"/>
    <f n="postpos" v="true"/>
    <f n="antepos" v="false"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>seul</form>
    <lem>seul</lem>
    <pos v="adv"/>
  </lex>
  <lex>
    <form>un</form>
    <lem>un</lem>
    <pos v="lex"/>

```

```

    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>un</form>
    <lem>un</lem>
    <pos v="det"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>un</form>
    <lem>un</lem>
    <pos v="det"/>
    <f n="subcat" v="ind"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>un</form>
    <lem>un</lem>
    <pos v="noun"/>
    <f n="subcat" v="num"/>
    <f n="gender" v="m"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>membre</form>
    <lem>membre</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>membr  </form>
    <lem>membr  </lem>
    <pos v="adj"/>
    <f n="postpos" v="false"/>
    <f n="antepos" v="false"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>

```

```

</lex>
<lex>
  <form>membre</form>
  <lem>membre</lem>
  <pos v="adj"/>
  <f n="postpos" v="true"/>
  <f n="antepos" v="false"/>
  <f n="number" v="s"/>
</lex>
<lex>
  <form>membre</form>
  <lem>membre</lem>
  <pos v="noun"/>
  <f n="subcat" v="abst"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>membre</form>
  <lem>membre</lem>
  <pos v="noun"/>
  <f n="subcat" v="conc"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>membre</form>
  <lem>membre</lem>
  <pos v="noun"/>
  <f n="subcat" v="hum"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>

```

```

<form>membre</form>
<lem>membre</lem>
<pos v="verb"/>
<f n="mode" v="ppast"/>
<f n="gender" v="m"/>
<f n="number" v="s"/>
</lex>
<lex>
<form>membre</form>
<lem>membre</lem>
<pos v="verb"/>
<f n="mode" v="ind|subj"/>
<f n="number" v="s"/>
<f n="pers" v="1|3"/>
</lex>
<lex>
<form>membre</form>
<lem>membre</lem>
<pos v="verb"/>
<f n="mode" v="imp"/>
<f n="number" v="s"/>
<f n="pers" v="2"/>
</lex>
<lex>
<form>de</form>
<lem>de</lem>
<pos v="lex"/>
<f n="case" v="min"/>
</lex>
<lex>
<form>de</form>
<lem>de</lem>
<pos v="XI"/>
</lex>
<lex>
<form>de</form>
<lem>de</lem>
<pos v="det"/>
<f n="subcat" v="ind"/>
</lex>

```

```

<lex>
  <form>dé</form>
  <lem>dé</lem>
  <pos v="noun"/>
  <f n="subcat" v="conc"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>de</form>
  <lem>de</lem>
  <pos v="prep"/>
</lex>
<lex>
  <form>l</form>
  <lem>l</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>
  <form>l</form>
  <lem>le</lem>
  <pos v="det"/>
  <f n="subcat" v="def"/>
  <f n="gender" v="f"/>
  <f n="number" v="s"/>
</lex>
<lex>
  <form>l</form>
  <lem>le</lem>
  <pos v="det"/>
  <f n="subcat" v="def"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
</lex>
<lex>
  <form>l</form>
  <lem>l</lem>

```

```

    <pos v="noun"/>
    <f n="gender" v="m"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>'</form>
    <lem>'</lem>
    <pos v="punc"/>
  </lex>
  <lex>
    <form>équipe</form>
    <lem>équipe</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>équipe</form>
    <lem>équipe</lem>
    <pos v="noun"/>
    <f n="subcat" v="conc"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>a</form>
    <lem>a</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>a</form>
    <lem>a</lem>
    <pos v="XI"/>
  </lex>
  <lex>
    <form>a</form>
    <lem>a</lem>

```

```

    <pos v="noun"/>
    <f n="gender" v="m"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>à</form>
    <lem>à</lem>
    <pos v="prep"/>
  </lex>
  <lex>
    <form>a</form>
    <lem>avoir</lem>
    <pos v="verb"/>
    <f n="mode" v="ind"/>
    <f n="number" v="s"/>
    <f n="pers" v="3"/>
  </lex>
  <lex>
    <form>été</form>
    <lem>été</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>été</form>
    <lem>été</lem>
    <pos v="noun"/>
    <f n="subcat" v="tps"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>été</form>
    <lem>être</lem>
    <pos v="verb"/>
    <f n="mode" v="ppast"/>
    <f n="gender" v="m"/>

```

```

    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>secouru</form>
    <lem>secouru</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>secouru</form>
    <lem>secourir</lem>
    <pos v="verb"/>
    <f n="mode" v="ppast"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>.</form>
    <lem>.</lem>
    <pos v="punc"/>
  </lex>
  <lex>
    <form>Il</form>
    <lem>il</lem>
    <pos v="lex"/>
    <f n="case" v="cap"/>
  </lex>
  <lex>
    <form>il</form>
    <lem>il</lem>
    <pos v="pro"/>
    <f n="procat" v="ppv"/>
    <f n="case" v="nom"/>
    <f n="gender" v="m"/>
    <f n="pers" v="3"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>entendu</form>
    <lem>entendu</lem>

```

```

    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>entendu</form>
    <lem>entendu</lem>
    <pos v="X"/>
  </lex>
  <lex>
    <form>entendu</form>
    <lem>entendu</lem>
    <pos v="adj"/>
    <f n="postpos" v="true"/>
    <f n="antepos" v="false"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>entendu</form>
    <lem>entendu</lem>
    <pos v="noun"/>
    <f n="subcat" v="hum"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>entendu</form>
    <lem>entendre</lem>
    <pos v="verb"/>
    <f n="mode" v="ppast"/>
    <f n="gender" v="m"/>
    <f n="number" v="s"/>
  </lex>
  <lex>
    <form>par</form>
    <lem>par</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>

```

```

</lex>
<lex>
  <form>par</form>
  <lem>par</lem>
  <pos v="noun"/>
  <f n="gender" v="m"/>
  <f n="number" v="s"/>
  <f n="proper" v="false"/>
  <f n="compound" v="false"/>
</lex>
<lex>
  <form>par</form>
  <lem>par</lem>
  <pos v="prep"/>
</lex>
<lex>
  <form>les</form>
  <lem>les</lem>
  <pos v="lex"/>
  <f n="case" v="min"/>
</lex>
<lex>
  <form>les</form>
  <lem>le</lem>
  <pos v="det"/>
  <f n="subcat" v="def"/>
  <f n="gender" v="f"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>les</form>
  <lem>le</lem>
  <pos v="det"/>
  <f n="subcat" v="def"/>
  <f n="gender" v="m"/>
  <f n="number" v="p"/>
</lex>
<lex>
  <form>lés</form>
  <lem>lé</lem>

```

```

    <pos v="noun"/>
    <f n="subcat" v="conc"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
  <lex>
    <form>lès</form>
    <lem>lès</lem>
    <pos v="prep"/>
  </lex>
  <lex>
    <form>les</form>
    <lem>le</lem>
    <pos v="pro"/>
    <f n="procat" v="ppv"/>
    <f n="case" v="acc"/>
    <f n="pers" v="3"/>
    <f n="number" v="p"/>
  </lex>
  <lex>
    <form>enquêteurs</form>
    <lem>enquêteurs</lem>
    <pos v="lex"/>
    <f n="case" v="min"/>
  </lex>
  <lex>
    <form>enquêteurs</form>
    <lem>enquêteur</lem>
    <pos v="adj"/>
    <f n="postpos" v="true"/>
    <f n="antepos" v="false"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
  </lex>
  <lex>
    <form>enquêteurs</form>
    <lem>enquêteur</lem>
    <pos v="noun"/>

```

```

    <f n="subcat" v="hum"/>
    <f n="gender" v="m"/>
    <f n="number" v="p"/>
    <f n="proper" v="false"/>
    <f n="compound" v="false"/>
  </lex>
</lexic>
<sfsa sz="19">
  <txt>Cinq marins sont toujours portés disparus,
    seul un membre de l'équipage a été secouru. </txt>
  <q id="0" pos="0">
    <tr lbl="0" to="1"/>
    <tr lbl="1" to="1"/>
    <tr lbl="2" to="1"/>
    <tr lbl="3" to="1"/>
  </q>
  <q id="1" pos="0">
    <tr lbl="4" to="2"/>
    <tr lbl="5" to="2"/>
    <tr lbl="6" to="2"/>
  </q>
  <q id="2" pos="0">
    <tr lbl="7" to="3"/>
    <tr lbl="8" to="3"/>
  </q>
  <q id="3" pos="0">
    <tr lbl="9" to="4"/>
    <tr lbl="10" to="4"/>
  </q>
  <q id="4" pos="0">
    <tr lbl="11" to="5"/>
    <tr lbl="12" to="5"/>
    <tr lbl="13" to="5"/>
    <tr lbl="14" to="5"/>
  </q>
  <q id="5" pos="0">
    <tr lbl="15" to="6"/>
    <tr lbl="16" to="6"/>
    <tr lbl="17" to="6"/>
    <tr lbl="18" to="6"/>
  </q>

```

```

    <tr lbl="19" to="6"/>
  </q>
  <q id="6" pos="0">
    <tr lbl="20" to="7"/>
  </q>
  <q id="7" pos="0">
    <tr lbl="21" to="8"/>
    <tr lbl="22" to="8"/>
    <tr lbl="23" to="8"/>
  </q>
  <q id="8" pos="0">
    <tr lbl="24" to="9"/>
    <tr lbl="25" to="9"/>
    <tr lbl="26" to="9"/>
    <tr lbl="27" to="9"/>
  </q>
  <q id="9" pos="0">
    <tr lbl="28" to="10"/>
    <tr lbl="29" to="10"/>
    <tr lbl="30" to="10"/>
    <tr lbl="31" to="10"/>
    <tr lbl="32" to="10"/>
    <tr lbl="33" to="10"/>
    <tr lbl="34" to="10"/>
    <tr lbl="35" to="10"/>
    <tr lbl="36" to="10"/>
  </q>
  <q id="10" pos="0">
    <tr lbl="37" to="11"/>
    <tr lbl="38" to="11"/>
    <tr lbl="39" to="11"/>
    <tr lbl="40" to="11"/>
    <tr lbl="41" to="11"/>
  </q>
  <q id="11" pos="0">
    <tr lbl="42" to="12"/>
    <tr lbl="43" to="12"/>
    <tr lbl="44" to="12"/>
    <tr lbl="45" to="12"/>
  </q>

```

```

<q id="12" pos="0">
  <tr lbl="46" to="13"/>
</q>
<q id="13" pos="0">
  <tr lbl="47" to="14"/>
  <tr lbl="48" to="14"/>
</q>
<q id="14" pos="0">
  <tr lbl="49" to="15"/>
  <tr lbl="50" to="15"/>
  <tr lbl="51" to="15"/>
  <tr lbl="52" to="15"/>
  <tr lbl="53" to="15"/>
</q>
<q id="15" pos="0">
  <tr lbl="54" to="16"/>
  <tr lbl="55" to="16"/>
  <tr lbl="56" to="16"/>
</q>
<q id="16" pos="0">
  <tr lbl="57" to="17"/>
  <tr lbl="58" to="17"/>
</q>
<q id="17" pos="0">
  <tr lbl="59" to="18"/>
</q>
<q id="18" pos="0" f="1"/>
</sfsa>
<sfsa sz="9">
<txt>Il a été entendu par les enquêteurs. </txt>
<q id="0" pos="0">
  <tr lbl="60" to="1"/>
  <tr lbl="61" to="1"/>
</q>
<q id="1" pos="0">
  <tr lbl="49" to="2"/>
  <tr lbl="50" to="2"/>
  <tr lbl="51" to="2"/>
  <tr lbl="52" to="2"/>
  <tr lbl="53" to="2"/>

```

```

</q>
<q id="2" pos="0">
  <tr lbl="54" to="3"/>
  <tr lbl="55" to="3"/>
  <tr lbl="56" to="3"/>
</q>
<q id="3" pos="0">
  <tr lbl="62" to="4"/>
  <tr lbl="63" to="4"/>
  <tr lbl="64" to="4"/>
  <tr lbl="65" to="4"/>
  <tr lbl="66" to="4"/>
</q>
<q id="4" pos="0">
  <tr lbl="67" to="5"/>
  <tr lbl="68" to="5"/>
  <tr lbl="69" to="5"/>
</q>
<q id="5" pos="0">
  <tr lbl="70" to="6"/>
  <tr lbl="71" to="6"/>
  <tr lbl="72" to="6"/>
  <tr lbl="73" to="6"/>
  <tr lbl="74" to="6"/>
  <tr lbl="75" to="6"/>
</q>
<q id="6" pos="0">
  <tr lbl="76" to="7"/>
  <tr lbl="77" to="7"/>
  <tr lbl="78" to="7"/>
</q>
<q id="7" pos="0">
  <tr lbl="59" to="8"/>
</q>
<q id="8" pos="0" f="1"/>
</sfsa>
</text-fsa>

```

