

Lucrarea nr. 5

Parsarea documentelor HTML

Mihai IVANOVICI

10 aprilie 2006

HyperText Markup Language (HTML) este limbajul cel mai folosit pentru a crea pagini web. Acest limbaj are ca scop realizarea unei modalități de prezentare către un client (browser) a informației existente pe un server.

Uneori se dorește extragerea de informații dintr-o pagină web și deci este nevoie de a procesa acea pagină web.

HTML este ierarhizat, ceea ce înseamnă că un parser trebuie să cunoască contextul în care apare un anumit “tag”. Acest lucru face ca metodele de parsing bazate pe expresii regulate să fie dificil de folosit. Parsingul devine și mai complicat atunci când paginile web nu respectă standardul HTML.

Python pune la dispoziție un modul `HTMLParser`. În această lucrare este ilustrată folosirea acestui modul.

1 Un exemplu simplu

Pentru a implementa un parser folosind modulul `HTMLParser`, se adaugă subclasei `HTMLParser` diverse funcții pentru tratarea diverselor tags. Să presupunem că vrem să parsăm următorul document HTML:

```
<!-- Basic title parsing example - basictitle.html -->
<HTML>
<HEAD>
<TITLE>My Title</TITLE>
</HEAD>
<BODY>
My text.
</BODY>
</HTML>
```

Următorul cod Python ilustrează modul de extragere a titlului din documentul de mai sus:

```
#Basic HTML Title Parser - basictitle.py
```

```

from HTMLParser import HTMLParser
import sys

class TitleParser( HTMLParser ):
    def __init__( self ):
        self.title = ''
        self.readingtitle = 0
        HTMLParser.__init__( self )

    def handle_starttag( self, tag, attrs ):
        if tag == 'title':
            self.readingtitle = 1

    def handle_data( self, data ):
        if self.readingtitle:
            #Title is usually small and simple
            self.title += data

    def handle_endtag( self, tag ):
        if tag == 'title':
            self.readingtitle = 0

    def gettitle( self ):
        return self.title

fd = open( sys.argv[ 1 ] )
tp = TitleParser()
tp.feed( fd.read() )
print "Title is: ", tp.gettitle()

```

Clasa `TitleParser` moștenește clasa `HTMLParser`. Metoda `feed()` a clasei `HTMLParser` va apela funcțiile `handle_starttag()`, `handle_data()` și `handle_endtag()`. Funcția `handle_data()` verifică doar dacă primește sau nu date de la un element `TITLE`, și în caz că primește salvează datele, concatenându-le.

Rulând exemplul de mai sus, veți obține:

```

$ python2.4 basictitle.py basictitle.html
Title is:  My Title

```

Exemplul poate fi extins pentru a extrage informația aflată sub formă de text în alte taguri (deschise și închise) din documentele HTML. De exemplu, pentru a extrage textul din tabele, veți extrage informația conținută de tagurile `<TR>` și `<TD>`.

2 Parsarea documentelor HTML

Exemplul precedent funcționează doar pentru documente HTML simple, corecte din punct de vedere sintactic. Nu toate documentele HTML sunt la fel de simple, și nu toate documentele cu extensia HTML conțin cod HTML valid. Browserele WEB curente sunt destul de îngăduitoare cu documentele HTML incorecte, și prin urmare acestea (documentele) au proliferat. Prin urmare, codul unei aplicații trebuie să ia în considerare aceste aspecte. Această secțiune prezintă diverse metode de a parse documente HTML, atât corecte cât și incorecte.

2.1 Translatarea entităților

În HTML se numesc *entități* caracterele obișnuite. De exemplu, entitatea HTML `&` reprezintă caracterul ampersant. Pentru afișarea sau parsarea codului HTML, de obicei se dorește convertirea entităților HTML în text pur. În continuare este prezentată o versiune modificată a primului exemplu, care ilustrează această problemă:

```
<!-- Entity title parsing example - etitle.html -->
<HTML>
<HEAD>
<TITLE>My Title &amp; Foreword</TITLE>
</HEAD>
<BODY>
My text.
</BODY>
</HTML>
```

Dacă veți rula codul precedent pentru acest document HTML, nu veți obține rezultatul dorit:

```
$ python2.4 basictitle.py etitle.html
Title is:  My Title  Foreword
```

Entitatea pur și simplu a fost ignorată, deoarece `HTMLParser` apelează funcția `handle_entityref()` atunci când întâlnește o entitate. Deoarece nu a fost definită această metoda, apelul funcției nu are nici un rezultat. Următorul exemplu ia în considerare entitățile:

```
#HTML Title Parser with Entity Support - etitle.py

from htmlentitydefs import entitydefs
from HTMLParser import HTMLParser
import sys
```

```

class TitleParser( HTMLParser ):
    def __init__( self ):
        self.title = ''
        self.readingtitle = 0
        HTMLParser.__init__( self )

    def handle_starttag( self, tag, attrs ):
        if tag == 'title':
            self.readingtitle = 1

    def handle_data( self, data ):
        if self.readingtitle:
            self.title += data

    def handle_endtag( self, tag ):
        if tag == 'title':
            self.readingtitle = 0

    def handle_entityref( self, name ):
        if entitydefs.has_key( name ):
            self.handle_data( entitydefs[name] )
        else:
            self.handle_data( '&' + name + ';' )

    def gettitle( self ):
        return self.title

fd = open( sys.argv[ 1 ] )
tp = TitleParser()
tp.feed( fd.read() )
print "Title is: ", tp.gettitle()

```

Rezultatul va fi de astă dată următorul:

```

$ python2.4 etitle.py etitle.html
Title is:  My Title & Foreword

```

Python pune la dispoziție mapări pentru entitățile HTML prin intermediul clasei `htmldefs`. Programul de mai sus utilizează codul entităților pentru a face conversia. Când o entitate este întâlnită, codul verifică dacă este recunoscută sau nu. Dacă da, atunci se face conversia, dacă nu, forma literală întâlnită în streamul de la intrare este afișată. Asta deoarece în documentele HTML deseori nu se folosește `&`, ci `&`, creându-se astfel entități incorecte.

2.2 Translatarea referințelor caracter

Pe lângă entități, documentele HTML pot conține referințe caracter. Acestea conțin valori întregi și arată de forma `®`. Referințele caracter sunt folosite pentru a fișă caractere care nu sunt întotdeauna tipăribile, cum ar fi caracterele din documentele HTML scrise într-o altă limbă decât engleza.

Iată un exemplu de fișier HTML ce conține o astfel de referință caracter:

```
<!-- Character reference title parsing example - ctitle.html -->
<HTML>
<HEAD>
<TITLE>My Title &amp; Foreword &#174;</TITLE>
</HEAD>
<BODY>
My text.
</BODY>
</HTML>
```

Codul funcției care tratează aceste referințe este următorul. Această funcție este adăugată ca metodă în clasa exemplului precedent.

```
def handle_charref( self, name ):
    #validate the name
    try:
        charnum = int( name )
    except ValueError:
        return

    if charnum < 1 or charnum > 255:
        return

    self.handle_data( chr( charnum ) )
```

Funcția va ignora referințele caracter invalide (cele care nu pot fi translatare într-un întreg sunt ^ în afara intervalului 0-255). Rulați exemplul `ctitle.py` având ca parametru de intrare fișierul `ctitle.html`. Observați rezultatul. Dacă nu vedeți un semn de “marcă înregistrată” la sfârșitul textului afișat, atunci s-ar putea ca terminalul să fie setat să afișeze alte caractere decât cele din setul Latin-1.

2.3 Tratarea tag-urile asimetrice

Cea mai supărătoare și des întâlnită problemă în codul HTML o constituie tag-urile asimetrice. De exemplu, o pagină de web ar putea avea un tag `<TITLE>` de început, dar să omită tagul `</TITLE>` de sfârșit.

Pentru unele tag-uri, HTML nu specifică obligativitatea închiderii lor, odată ce au fost deschise, cum ar fi <P> și LI. De aceea trebuie create mecanisme de tratare a acestor probleme. Standardul XHTML obligă închiderea tuturor tag-urilor.

În continuare este prezentat un exemplu de document HTML cu tag-uri asimetrice:

```
<!-- Unbalanced tags parsing example - utitle.html -->
<HTML>
<HEAD>
<TITLE>My Title &amp; Foreword &#174;
</HEAD>
<BODY>
My text.
<UL>
<LI> Primul element din lista,
<LI> Al doilea element din lista,
<LI> Al treilea element din lista.
</BODY>
</HTML>
```

Veți observa că tag-ul </TITLE> de sfârșit lipsește, la fel și tag-ul . Codul următor va trata aceste situații:

```
#HTML Extended Parser - utitle.py
```

```
from htmlentitydefs import entitydefs
from HTMLParser import HTMLParser
import sys, re
```

```
class TitleParser( HTMLParser ):
    def __init__( self ):
        #the taglevels list keeps trak of where
        #we are in the tag hierarchy
        self.taglevels = []
        self.handletags = [ 'title', 'ul', 'li' ]
        self.processing = None
        HTMLParser.__init__( self )

    def handle_starttag( self, tag, attrs ):
        """Called whenever a start tag is encountered"""
        if len( self.taglevels ) and self.taglevels[-1] == tag:
            #processing a previous version of this tag; close it out
```

```

        #and then start a new one on this one.
        self.handle_endtag( tag )

#note that we're now processing this tag
self.taglevels.append( tag )

if tag in self.handletags:
    #only bother saving off the data if it's a tag we handle
    self.data = ''
    self.processing = tag
    if tag == 'ul':
        print "List started."

def handle_data( self, data ):
    """This function simply records incoming data if we are
    presently inside a handled tag."""
    if self.processing:
        #this could be slow for large files; for this example
        #it's a simple way to save off data.
        self.data += data

def handle_endtag( self, tag ):
    if not tag in self.taglevels:
        #we didn't have a start tag for this anyway
        #just ignore it
        return

    while len( self.taglevels ):
        #obtain the last tag on the list and remove it
        starttag = self.taglevels.pop()

        #finish processing it.
        if starttag in self.handletags:
            self.finishprocessing( starttag )

        #if it's our tag, stop now.
        if starttag == tag:
            break

def cleanse( self ):
    """Removes extra whitespace from the document."""
    self.data = re.sub( '\s+', ' ', self.data )

def finishprocessing( self, tag ):

```

```

        self.cleanse()
        if tag == 'title' and tag == self.processing:
            print "Document title:", self.data
        elif tag == 'ul':
            print "List ended."
        elif tag == 'li' and tag == self.processing:
            print "List item:", self.data

        self.processing = None

    def handle_entityref( self, name ):
        if entitydefs.has_key( name ):
            self.handle_data( entitydefs[name] )
        else:
            self.handle_data( '&' + name + ';' )

    def handle_charref( self, name ):
        #validate the name
        try:
            charnum = int( name )
        except ValueError:
            return

        if charnum < 1 or charnum > 255:
            return

        self.handle_data( chr( charnum ) )

fd = open( sys.argv[ 1 ] )
tp = TitleParser()
tp.feed( fd.read() )

```

Rezultatul va fi următorul:

```

$ python2.4 utitle.py utitle.html
Document title: My Title & Foreword
List started.
List item:  Primul element din lista,
List item:  Al doilea element din lista,
List item:  Al treilea element din lista.
List ended.

```

În funcția `handle_starttag()`, în `self.taglevels` se salvează fiecare tag întâlnit. Dacă tag-ul este unul dintre cei trei pe care programul îi

tratează, atunci se setează flag-ul `self.processing` pentru a semnala faptul că s-a început salvarea datelor.

În funcția `handle_endtag()`, se verifică dacă a fost întâlnit un tag de început, pentru cel de sfârșit întâlnit. Dacă nu, se ignoră tag-ul respectiv. Dacă tag-ul de început a fost întâlnit, atunci se va găsi cea mai recentă apariție, prin parcurgerea în sens invers a listei `self.taglevels`, începând cu sfârșitul acesteia.

Funcția `finishprocessing()` începe prin a elimina spațiile albe din șirul de caractere de la intrare. Verificarea `tag == self.processing` are rolul de a asigura că aceleași date nu sunt folosite de două ori. De exemplu, dacă tagurile de interes ar fi fost la trei nivele adâncime, exemplul de mai sus ar fi salvat doar datele de pe nivelul cel mai recent întâlnit.