

We are IntechOpen, the world's leading publisher of Open Access books Built by scientists, for scientists

6,900

Open access books available

185,000

International authors and editors

200M

Downloads

Our authors are among the

154

Countries delivered to

TOP 1%

most cited scientists

12.2%

Contributors from top 500 universities



WEB OF SCIENCE™

Selection of our books indexed in the Book Citation Index
in Web of Science™ Core Collection (BKCI)

Interested in publishing with us?
Contact book.department@intechopen.com

Numbers displayed above are based on latest data collected.
For more information visit www.intechopen.com



An Implementation of Intelligent HTMLtoVoiceXML Conversion Agent for Text Disabilities

Young Gun Jang
Chongju University
Republic of Korea

1. Introduction

HTML represents information in a visual manner and as a result text disabilities such as visually impaired individuals and dyslexics are unable to access their content. Beyond doubt the vast majority of modern web applications neglect the special needs of disabled people. The acoustic representation of the web content contributes to the purpose of pervasive learning and offers great help to disabled individuals. Proposed application of this paper is based on VoiceXML in order to become accessible acoustically by a standard phone or by a computer.

Almost all internet contents are constructed by HTML and the contents which internet user prefer are various and extensive. Manual contents conversion from HTML to VoiceXML for the general public needs so much costs and seems to be practically impossible. Therefore automatic or semiautomatic conversion is very necessary for practical acoustic service.

The possible applications on the internet and intranet are as various as the contents contained on the internet and as the users who utilize the information. This variety increases the attempts of making an automated internet for its complexity of information and service. Especially, these attempts are very necessary for the web contents processor to decrease the cost of producing and managing web contents which is increasing rapidly, also to minimize the number of mistakes caused by human error. Including filling out forms, web automation should automatically operate the web activities, and it should also operate according to the plans or requirements. However, information on the internet often changes its contents, and even its structure. Because HTML describes more for visual expression than for the contents or structure, web pages are still unable to be interpreted perfectly by present technologies (Asakawa, 2000).

VoiceXML is an eXtensible Markup Language (XML) based language that aims to function as a tool for the development of interactive voice response applications. VoiceXML is designed for creating voice applications that feature synthesized speech, audio recognition of voice input or Dual-Tone Multi-Frequency input, recording of spoken input, reproduction of audio files, control of dialog flow, and telephony features, such as call transfer. In general terms VoiceXML attempts to constitute the audio equivalent of

Extensible Hypertext Markup Language (XHTML) for the voice web instead of the visual web.

Before VoiceXML, VoxML which is one of the original types was published by Motorola, also Goose, et al was the first to convert HTML into VoxML. Goose, et al presented the study of Vox portal which makes it possible to access the web through a phone-line, also has 4 layer structures such as client, agent, web server, database, which was added to the VoxML-Agent that can convert HTML to VoxML in the traditional 3 layer WWW structure. Concerning the VoxML-Agent, the interaction with a Vox portal was mainly mentioned, however, the core part which is the design of the conversion function was not mentioned (Goose et al., 2000).

Building a vocal interface by interpreting HTML documents, which use mainly visual interface, is currently being studied (Mohan et al, 1999; Asakawa et al., 1998; Vankayala et al., 2006), also code conversion based on the grammar characteristics is being studied mainly with a tool to adjust and access the internet contents (Embley et al., 1999), code conversion based on the structure characteristics of HTML (Goble et al., 2001; Huang, 2000; Buttler, 2001; Choi, 2001), Semantic code conversion(Hori et al., 2000; Krulwich, 1997), manually adding annotation code (Asakawa, 2000; Lieberman et al., 1999), code conversion used structure and format information (W3C, 2000), semi automatic code conversion using the interaction between structure model and computer (Li et al., 2004), however, these studies are not being generalized because the purpose of the study is to rearrange web pages and convert web objects for the easier voyage and understanding and only to solve a specific user group's needs. It is impossible for not only computer but also humans who know HTML well to guess visual layouts, and to separate each content perfectly just by reading HTML code. Therefore, it is necessary to limit the converting subject and intelligent methods of selecting, extracting and separating the subjects. The converter selects the contents, separates and extracts selected contents, and converts the extracted contents into a VoiceXML document according to the priory defined scenario. The core technology of the converter is to select, and extract the contents that the users want, regardless of the expression of HTML and to accurately separate the contents according to the statistical, grammatical, and structural characteristics of an HTML document and to increase the automatic rate of converting contents by one time access for the practical purpose. Regarding the separation of contents, Embley (1999,2006), Buttler (2001) and Choi (2001) suggested a method of separating multi contents that are arranged in a row with the similar contents like a bulletin board, according to the contents in the HTML document, that has multi contents, and also suggested Heuristic algorithms that can extract separation tags which is the standards of the separation by using the structural characteristics of an HTML document. For the limitation of the suggested Heuristic algorithm, Choi (2001) limited the subjects of the multi contents type to bulletin boards, lists and searching results. The limited subjects should be selected safely on the web page for the successful conversion. However, in the case of using the structural characteristics to select the contents, the contents can be selected differently according to the various expressions and changes of contents on a web page. The minimum child node that is used by Embly and Goose can't properly reflect the user's intention, also the minimum child node that is suggested by Choi and node selection methods which combines included characters and numbers can make different selections according to the number of nodes and the balance of included numbers of characters. Therefore, other methods are required to use the structural characteristics. In the semantic

approach, it suggests using ontology for this problem (Huang, 2000; Goble et al., 2001). The documents should be described on an XML base because it is based on the meaning, also it requires an RDF to interpret the ontology. However, it is very complicated to describe the ontology to select contents on the web page that we already know, and it requires another RDF and professional knowledge about the technical language and grammar, moreover, if it doesn't exactly match with the ontological specifications, it responds to the slight changes of the contents very sensitively (Goble et al., 2001). To reflect the user's intention, Krulwich (1997) used the method in which the agent records the interaction between the user and the web page so that, when the user connects to the web page again, the agent performs the same work that is selected by the user. However, this method is very difficult to implement, also it uses the HTML location defining language as the converted result during the automatic process. Lieberman of MIT suggested a way to make the agent recognize the characters through training (Lieberman et al., 1999). This method removed the professionalism that is necessary while manually entering the grammar and eliminating the possibility for mistakes by suggesting examples to the agent and by training the semantic character patterns which are in the unstructured web information. These studies concerning the selection of contents all have some advantages, however, they are not very convenient and strong because they might require unnecessary information to be defined by the users or provide unnecessary interface, also they are very complicated to implement.

This study suggests; 1) the hybrid sequential contents group selection method that utilizes a variety of factors including structural characteristics, the users' prior knowledge, interaction with the agent, and character information to ensure the inclusiveness and definiteness of the contents based access methods to solve the problem of selecting contents accurately, and it also maintains advantages such as there is no need for professionalism and convenient character recognition and it covers the disadvantages. 2) the sequential recognition for documents based on schematics that provide documents interactively according to the user's prior knowledge regarding the relationship between the documents, which is to increase automation rate of the converter. 3) converting the entire structure of HTML to VoiceXML, which is applied the above suggested 1) and 2), according to the content separation which is the basic system and designated scenario. In Section 2, it is suggested the document type for conversion, in Section 3 and 4, it is suggested the contents selection method and recognition methods of sequential web documents and in Section 5, it is suggested the converting from HTML to VoiceXML agent structure which is combined the suggested methods, and constructing methods, also described the contents separation method.

2. Selection by the convertible HTML document type and structural characteristics

A web author tries to put a lot of information on one page by using various visual effects. This information is visually categorized and then is fragmented, For example, a web page in Fig. 1 has A, B, C, D, E, F, G and H fragments on one page. It is necessary to determine which fragment is the main information on this page and to separate the information according to the contents and then it should be converted to VoiceXML document (W3C, 2000).

In order to convert HTML to VoiceXML, first of all, the HTML contents should be separated into fragment units according to the contents, then a VoiceXML document should be created

for each fragment. However, HTML tags are difficult to separate into fragments because they only express how to visualize the information, and don't include any meaning as an XML tag.

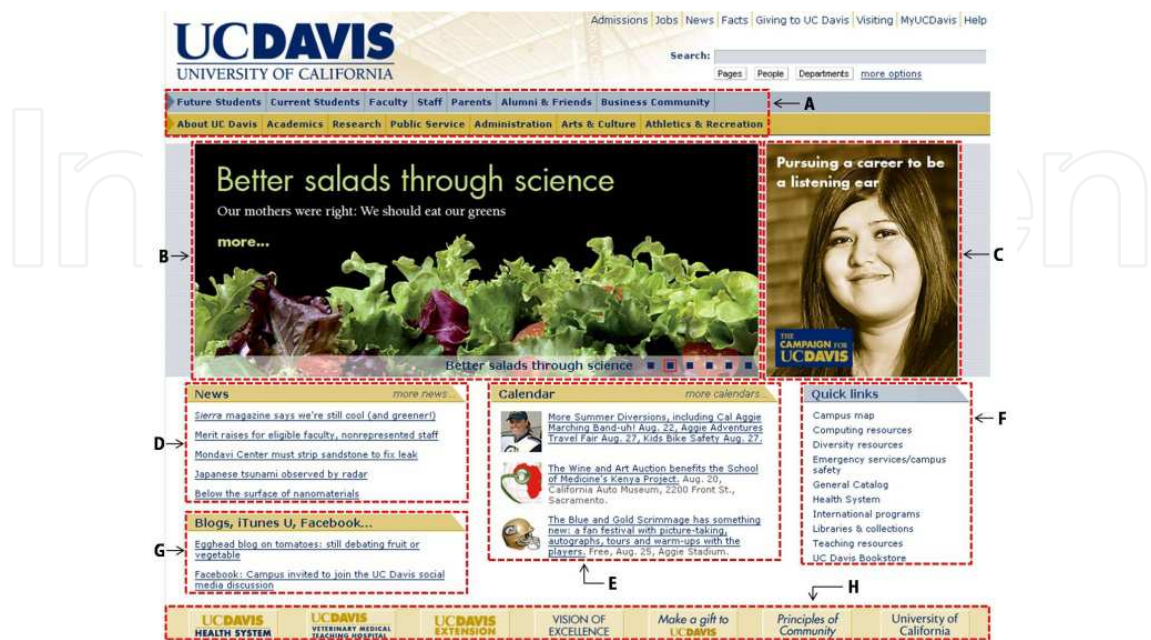


Fig. 1. Web page type.

Therefore, it is assumed that the contents are visually categorized. The probability of this assumption was proved through the code verification of most web pages. For the detailed separation of the contents into fragments, the heuristic method was used, which statistically analyzed the structural characteristic of the tags. Therefore, the types that can be separated by using this method are limited, also the qualitatively similar types of contents are arranged in a row. The documents types that can be applied the conversion method of this study are bulletin board, list and search result. This type of HTML document has many features that are similar to child nodes when the documents are expressed in a tree structure, also the contents include many characters, which means the location of the contents can be found. The contents that can be converted to VoiceXML in Fig. 1 are the News(D), Calendar(E) and Blogs, iTunes U, Facebook(G).

To analyze the structure of an HTML document, for the first step, it must be reorganized into a tree structure. The reason is because it is easier to analyze an HTML document in a tree structure.

The minimum reformed sub tree method is designed as the selection method of the main content group by using the structural characteristics on the web page. The minimum sub tree means the smallest tree that includes contents of the whole tree structure. Before the contents are extracted, it first extracts the sub tree that includes the contents, and this sub tree corresponds to each fragment of Figure 1. The minimum sub tree is the tree that uses the nodes with the most child nodes as the path after finds out how many child nodes are in each node. Usually, in the web page where contents are arranged in a row, the possibility that the main contents are included in the sub tree that uses the node with the most child nodes as its path, is very high. However, if the numbers of child nodes becomes the main

concern of extracting the minimum sub tree, in the web page which has lots of menus, the minimum sub tree can be the sub tree that includes the menu.

In this paper, it is suggested that the text size of the node with the most child node be used to extract the minimum sub tree. Meaning that first it is necessary to find the node with the most child nodes, then figure out the text size of each node and finally the minimum sub tree should be the sub tree that uses the node with the biggest text size as its path. In this method, we can avoid the mistake of designating the sub tree that has most child nodes with the small text size, like a menu, as its minimum sub tree.

3. Selecting an interactive contents group based on contents

When selecting a contents group through the structural interpretation of an HTML document, it is difficult to apply the user's intention with any kind of rule platform, and difficult to select multiple contents on one web page. Therefore, there must be a mechanism to select contents again by reflecting user's intention if the automatically selected contents by using the structural characteristics are not reflecting user's intention. The suggested method is the sequential selection method which provides the automatic selection and the user's multi selection according to the frequency of use based on the interactive selection with users. It first applies the structural method to select the proper item, and if it is improper, then it can be selected suggesting examples through the character information and by checking the composed results on the web page.

At this time, the suggested character row becomes the title of the documents that are created as the result of the conversion. Also if the wrong contents group is selected, it can be selected by designating numbers or numbers with logical calculus. In this paper, character rows included in a web page were used to train the agent. Web authors usually try to show the contents to users in a simple and effective way by using key words as characters or graphics. Therefore, in this method, it uses the author's expression ability and insight about the web page that is showed to the users, and copy key words are selected for the contents to transfer to an agent, also it designates a sub tree that includes the valid key words among the already written sub tree as the converting subject for the structural interpretation. However, it is unable to copy the characters of key words if they are expressed graphically, and this method can't be used if multiple contents or logically composed contents need to be selected. Contents composition looks like the wrong selection or a visually same contents group, however, it can be used very effectively on web documents that are made with a different structure. If there is a problem for selection from the character row, it can include the index number of contents, which are made according to the result of structure interpretation, on the original document, and can visually provide the composed web document to the user, also it can select the contents by transferring the index number to the agent according to user's judgments. Multiple selections is possible in this method, also it supports logical composition + operator which means combining is only used for the logical composition operator, however, it can be added by needs.

Fig. 2 is an example of the valid implementation algorithm. In this method, the structure interpretation method as tree structure combined with the rule platform that uses the amounts of nodes and characters, also if the automatic selection is failed because it doesn't reflect user's intention: there are many web making methods that can have a similar visual

effect: it is impossible to use for all methods that can be used on HTML interpretation, so it utilizes a hybrid sequential method with a meaning based access that allows users to interactively select the interpreted results of the system.

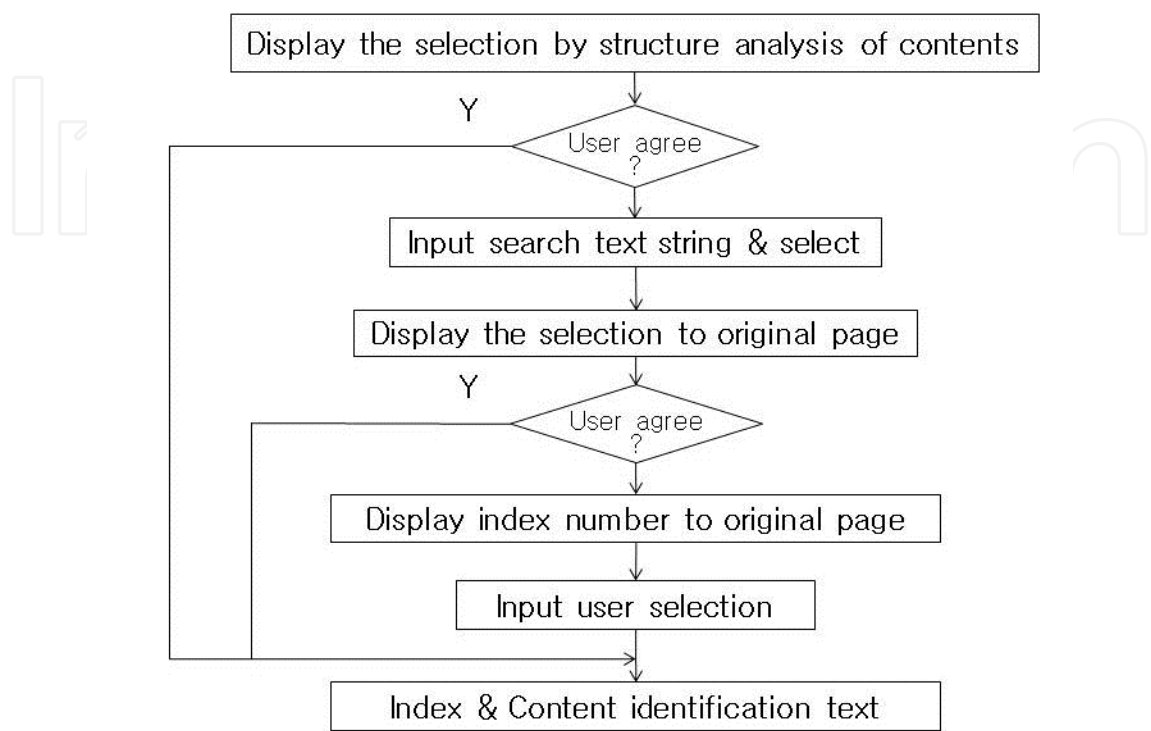


Fig. 2. Flow of the hybrid sequential contents group selection method.

4. Sequential recognition of web document based on interactive schema utilizing the prior knowledge

Generally, if there are many contents that have the same characteristics and structure, the contents should be made with several web documents and those web documents should be connected through the link because the amount of contents that can be expressed on one web page is limited. However, it is unable to determine which link is about the next document by only looking at the HTML code, so it is impossible to extract automatically at once. To solve the above problem, the entire contents can be extracted by interactively providing the user's prior knowledge regarding the link between the documents, and this information doesn't designate the tag path by analyzing the HTML code, but provides the link information like the link is located after how many links from the last contents so the URL for the next document can be extracted. At this time, the link between documents is usually located after the last contents, and the link that connects the next document can be divided into three categories; text, image and serial number, as seen in Fig 3.

The prior knowledge information input scheme can be divided into 3 steps as seen in Fig. 4. In the first step, the link type for the next document is designated, and in the second step, the details such as link character, link location and serial number orders is designated according to the link type, and in the last step, the location of the last document is designated, like how many documents will be extracted.

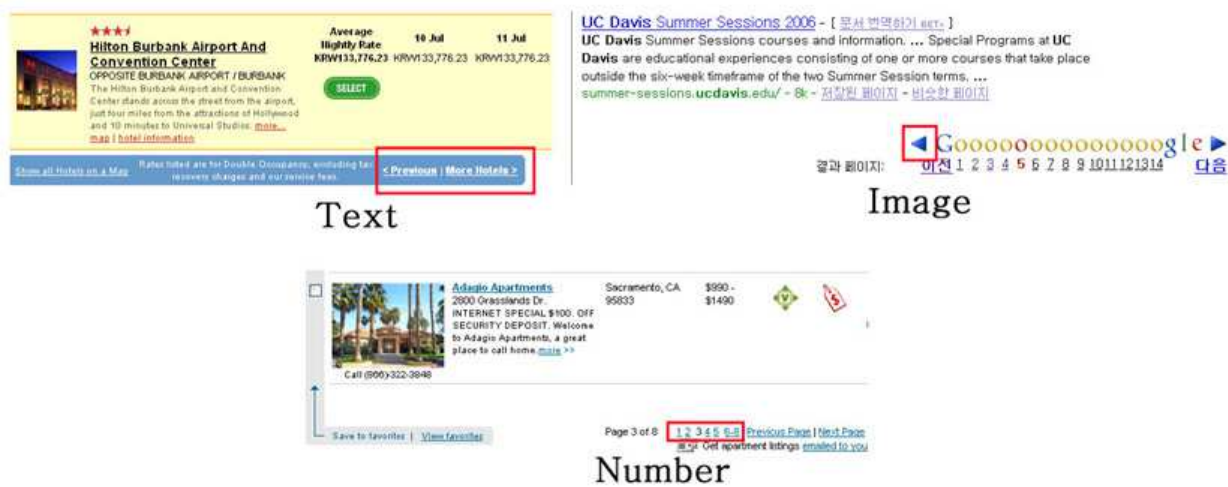


Fig. 3. Linking type for the connected page and link location.

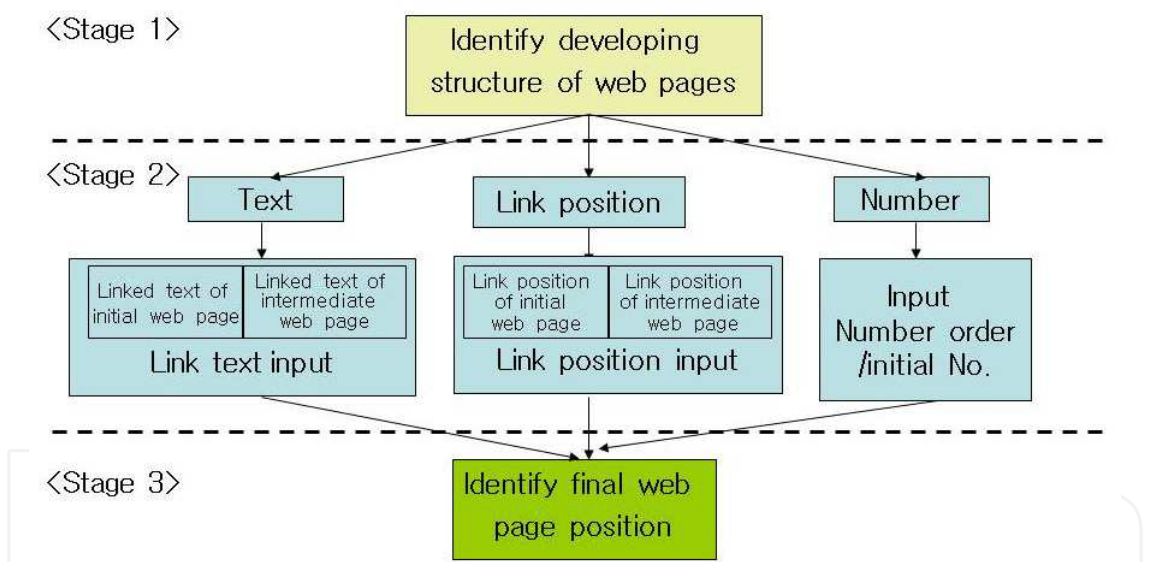


Fig. 4. Prior Knowledge Information Input Schema.

When the prior knowledge is provided, the URL for the next linked document can be extracted through the procedure in Fig. 5.

The prior knowledge information input procedure for the HTML document structure is provided in a Windows Wizard style. The Prior knowledge information input wizard can be divided into 3 steps as the Prior knowledge information input schema. In the 1st step, the link type information should be entered, then, in the 2nd step, if the link type is text, enter the text information, if it is image, enter link location information, and if it is serial number, enter serial information. In the 3rd step, enter the final web document information which is needed to be extracted.

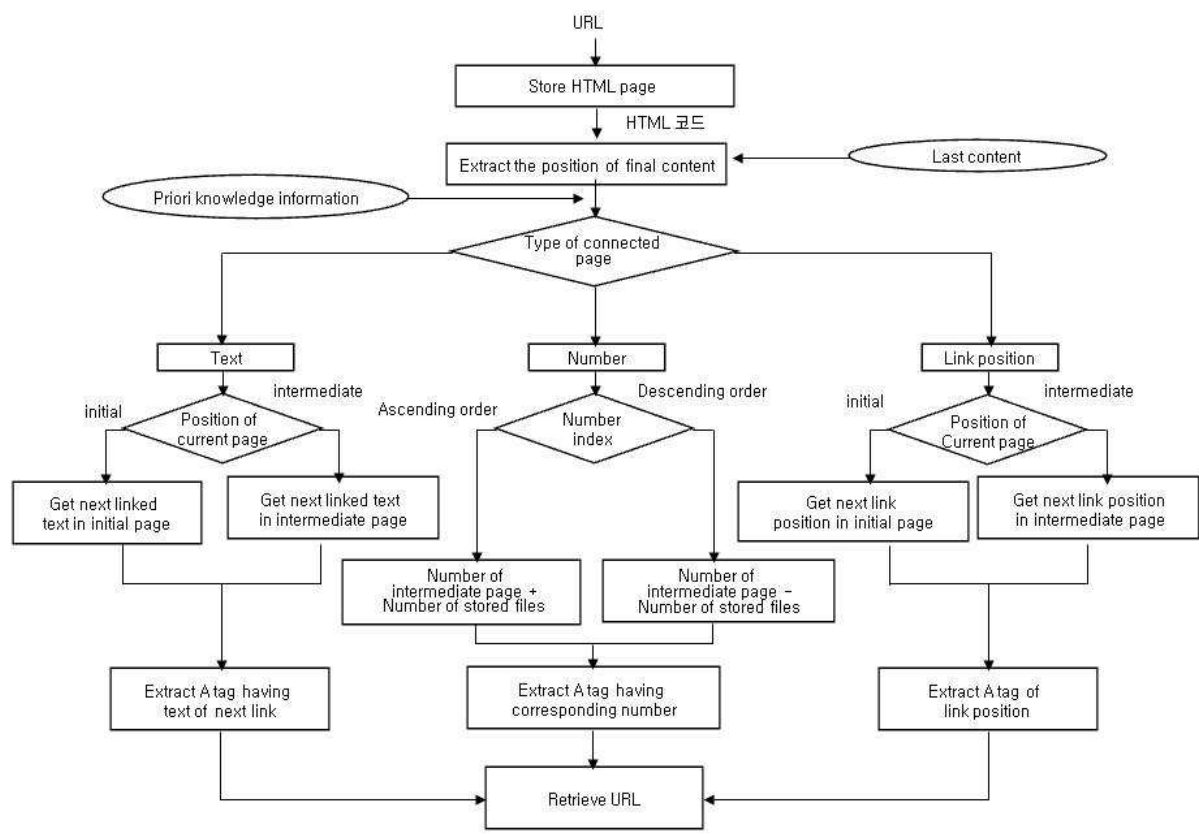


Fig. 5. Extracting URL Flow for the connected web documents.

5. Design & implement HTML to VoiceXML conversion module

The simplest way to change HTML into a markup document is to match the tags with similar roles with another markup language to the tags in the HTML. This method of conversion is usually applied when we convert an HTML document into a markup document for wireless internet. However, HTML tags don't provide any meaning but the visual expression of the contents, so it is difficult to convert and provide exact information by using 1:1 matching for each tag. In another method, we can reorganize the scenario by recording the web documents path and find the contents using that path so the information can be provided with the markup language. At this time, the web document path will be recorded as the absolute path that uses html or a body tag as its path in the tree structured HTML(Freire et al, 2001). The problem of this method is that it can cause unexpected results because the tags were deleted, modified, and inserted for the path. Embley, Butler and Choi suggested separating the multiple contents that are arranged into similar types such as bulletin board or searched results by using the documental structure, however, it is only about one web document, so we need to connect each web document to separate the sequential contents across several web documents.

The HTML to VoiceXML Conversion Agent that is implemented in this paper can be divided into three procedures as seen in Fig. 6; Extracting list and creating VoiceXML

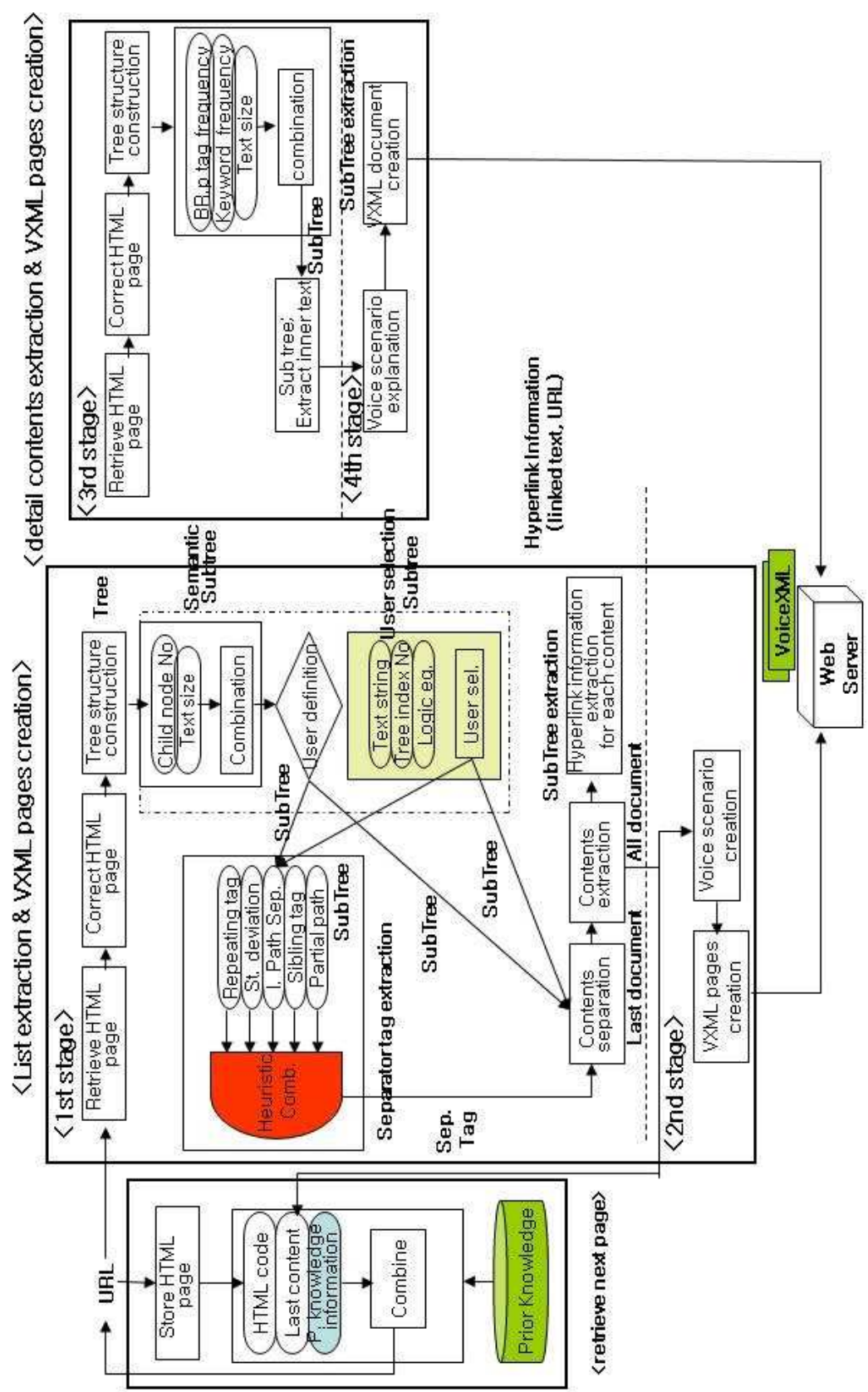


Fig. 6. HTMLtoVoiceXML Conversion Module Map.

document, Extracting details and creating VoiceXML document, and Extracting URL of the next connected web document. Extracting list and creating VoiceXML document procedure is extracting list contents from the list HTML document and then creating the VoiceXML document that is valid for the list scenario. The dotted line part is selected contents after the conversion in the tree structure in the first step, and it is implemented through the hybrid sequential contents selection algorithm that is suggested in this paper. In this step, the contents selection method which has faster processing speed and suggested in section 2, will be applied, also if it is not the contents wanted, it will be applied the method that is suggested in section 4. In the process of extracting details and creating a VoiceXML document, it extracts the details from the detail contents document about the list contents and creates it with VoiceXML documents which will be valid for the detailed contents scenario. In the process of extracting the URL from the next connected web document, if the contents are across several web documents, it extracts the list and creates the VoiceXML document after finding out the URL of the connected web page. It finds the URL of the connected web document through prior knowledge of the list HTML document structure, extracts the list, and creates the VoiceXML document to extract all the contents during a single connecting by transferring the value. This part has not been suggested in other existing studies, however, it extracts connected web documents by using the connection rule which determined by the connected web document recognition method, then transferring that information to the 1st step and continuing to perform the same procedure. The dotted line part which is the content selection block and the solid line part which is extracting URL block for the connected web page in Fig. 6 opposite with the data block that uses the first interpretation result during the repeated connecting procedure, also it is connected continuously by the scheduler.

First of all, the candidate tags for the separation of contents should be extracted in order to separate contents to each unit. The contents candidate tags should be the child nodes of the path nodes of the selected sub tree.

If there is only one candidate tag for separating the contents, that candidate tag will be the boundary tag to separate the contents, however, if there are more than two candidate tags, then the valid tag should be selected as the separation tag. To extract a valid separation tag from several separation candidate tags, the heuristic algorithm was used, which is based on a few rules used when making HTML document. The heuristic that is used in this paper makes: standard deviations about the amounts of characters between candidate tags: repeating patterns of the tag pair which consistently shows on each contents: all path lists from candidate tag nodes to other temporary nodes then, counts how many times the partial paths that are appeared by the identified path happened and sibling tags that adjoin to sub tree happened, then, rearrange the sibling tags that categorize all tags in descending order and tags which are frequently used to separate contents in the ascending order, then finally use the identifiable path separator tag (IPS) for the candidate tags. Each heuristic is optimized for the special type of web page, but it is independent from other heuristic algorithms. Therefore, it is considered to combine each independent heuristic to increase the possibilities of extracting correct contents separation tags of web document.

To decide on the best combination status for 5 heuristic algorithms, the Stanford Certainty Theory was used (Luger et al., 1997). There are a total of 16 ways to combine each heuristic algorithm, however, most the exact way was to extract the separation tags after combining 5 heuristic algorithms.

When the contents are separated for each separation tag, the starting and ending points of the contents should be extracted. Meaning that it is necessary to first decide if the starting tag of the separation tag and the next starting tag of the separation tag should be the starting and ending point, or if the contents right before the separation tag and the starting tag of the right next separation tag should be the starting and ending point to separate contents. Separating procedure with the standard of the separator tag, which is suggested in this paper, is as following.

```

1. Find the first child node N of the path node of the minimum sub tree
2. repeat
  if (N == STYLE Tag or SCRIPT Tag or !(Annotation) Tag)
  or ( N == BR Tag and BR Tag != Separation Tag)
  or ( N == HR Tag and HR Tag != Separation Tag)
  or ( N == P Tag and P Tag != Separation Tag)
  then
    N := Next siblingnode
    continue /* Perform the if sentence again */
  else
    End the repeat sentence
  end if
until end-of-Minimum sub tree
3. if N == Separation tag then
  Contents := The contents between the starting tag of the first separation tag
  and the starting tag of the next separation tag
else
  Contents := The contents between the starting tag of the first separation tag
  and the starting tag of the next separation tag
end if

```

When the contents are separated with the standard of the separator tag, the invalid contents can be included. Therefore, the valid contents should be extracted from the separated contents. In each text of contents of the web page, there is an arrangement of regular types of numbers, dates and characters, so it is necessary to use these characteristics to extract the valid contents. The suggested method of extracting valid contents in this paper is as follows:

```

1. repeat /* For all separated contents */
  Text type of the separated contents should be divided into numbers, dates
  and characters to be saved at F arrangement
until end-of-Separated contents
2. Find the amounts of the same types from the saved information at F
arrangement.
3. Find the most types of information and extract the contents that have this
type of information.

```

The extracted contents through the extracting procedure shall be provided to the users through the vocal interface. Because users can understand only a limited amount of vocal information at once, compared to visual information, N vocal scenarios shall be created according to the amount of extracted contents.

In this paper, if the average characters of the extracted contents is less than 100, a vocal scenario with 4 contents will be created, and if it is more than 100, N scenarios will be created consisting of 2 contents in one vocal scenario.

Each vocal scenario creates VoiceXML document according to the VoiceXML 1.0 format. The created VoiceXML document has the basic structure displayed in Fig. 7.

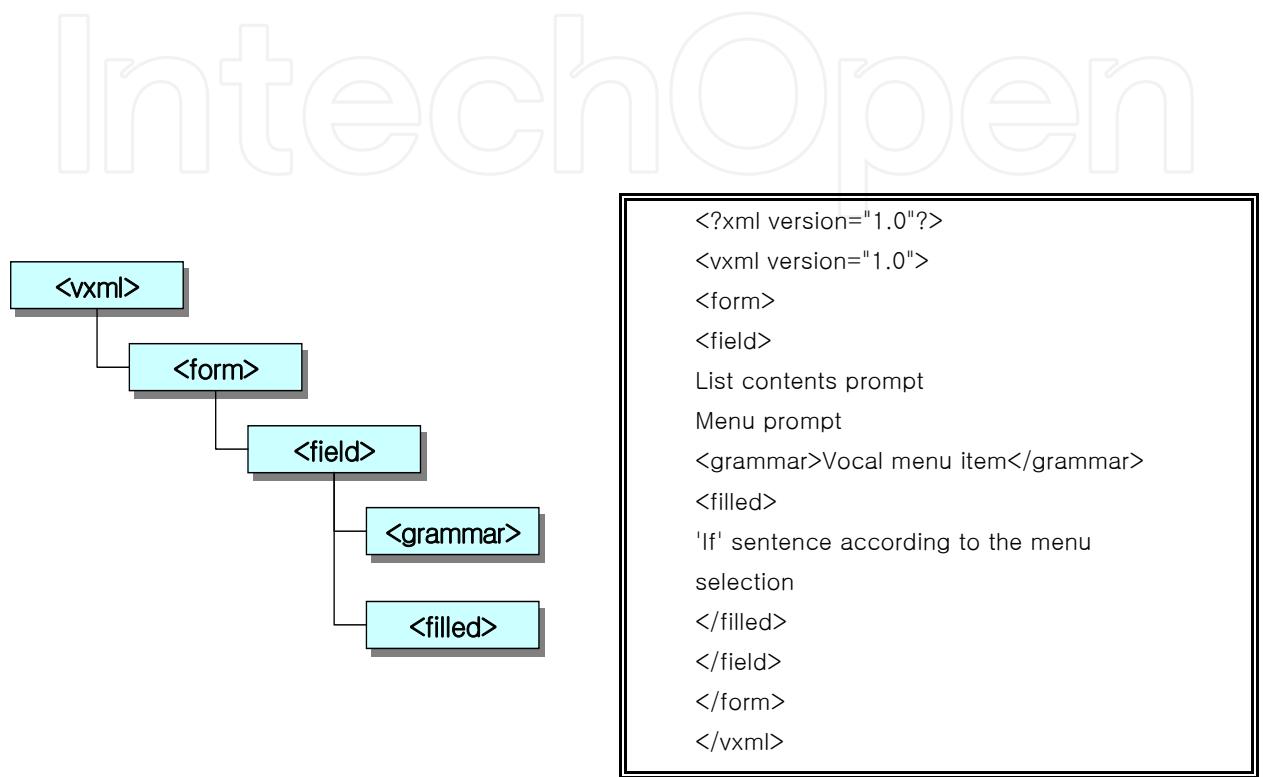


Fig. 7. Basic structure of the created VoiceXML document.

6. Test & results

To test the performance of the HTML to VoiceXML conversion agent that was implemented in this paper, board type, list type and search result type employment sites, newspaper sites and search engines with high access frequency were selected. The specific sites selected for test are listed in table 1, and the sites that were tested for all 3 types are listed in table 2.

The test method is the same as the method in GIT, and the succeed rate of algorithm will be counted in two steps (Li et al., 2004). First, the suggested algorithm was applied to each web site page, the tag that with the highest score was found and the exact rate of the tag on that page was counted. Second, an average of all web sites was calculated to see the success rate of the heuristic algorithm and the combined algorithm.

In the result of applying this method to 200 web pages, which are shown on table 1, the separation successful rate was 100% except for the pages that described the wrong HTML grammar. This was because most of the web sites consisted of tables, and in many case, it had only one separation candidate tag, so the heuristic algorithm was not needed.

Division	Name of Web site	Web site URL
Employment (Board type)	Incruit	www.incruit.com
	Recruit	www.recruit.co.kr
	Hellojob	www.hellojob.com
	Jobex	www.jobex.co.kr
	Devpia	www.devpia.com
Daily News (List type)	DongA Daily	www.donga.com
	JoongAng Daily	www.join.com
	Sports Today	www.sportstoday.co.kr
	GoodNews	www.kukminilbo.co.kr
	Hankyoreh	www.hani.co.kr
Search (Search results type)	MSDN	msdn.microsoft.com
	Paran	www.paran.com
	Altavista	www.altavista.co.kr
	Google	www.google.com
	Naver	www.naver.com

Table 1. Test subject web sites.

Name of Web site	Web site URL
Hankooki	www.hankooki.com
Hankook Ilbo	www.hankooki.com/hankook.htm
Daily Sports	www.hankooki.com/dailysports.htm
Seoul Economy	www.hankooki.com/sed.htm
Korea Times	www.hankooki.com/sed.htm

Table 2. Test subject web sites for success rate.

Heuristic Algorithm	No. 1	2	3	4
Separative Tag	0.85	0.1	0.15	0.0
Repeating Pattern	0.65	0.0	0.15	0.0
Partial Tag	0.95	0.0	0.0	0.0
siblingTag	0.80	0.2	0.0	0.0
Standard Deviation	0.70	0.2	0.1	0.0
Combination Algorithm	0.99	0.02	0.01	0.0

Table 3. The probability ranking of each algorithm and the success rate of the combined algorithm.

Therefore, we tested the success of the separation tag extraction rate for the web sites that used various web page writing methods. In table 3, the probability ranking of each heuristic algorithm and the success rate of the combined heuristic algorithm is described for about 200 web pages from the web sites suggested in table 2. The reason it had a higher succeed rate compared to the GIT combination algorithm success rate of 94%, is because the different types of web page were less than GIT. It is difficult to compared with the GIT case, because there is no information about the web page. However, the test results show that the implemented converter is very trustworthy for the regular table or the contents that links are arranged in a row, also when the combined heuristic algorithm was applied for the web page that used various methods, it showed a 99% success rate, so most of the web pages were convertible. In table 3, it shows the success rate when each heuristic algorithm is applied for the separation candidate tags of the web pages, and when each heuristic is combined.

Web Site Name	Prior Knowledge Type
Incruit	Numbers, Image
Recruit	Numbers, Image
Helljob	Numbers, Image
Jobex	Numbers, Text
Devpia	Numbers
DongA Daily	Image
JoongAng Daily	Image
GoodNews	Image
Hankyereh	Numbers
MSDN	Numbers
Paran	Numbers
Altavista	Numbers
Google	Numbers, Text
Naver	Numbers

Table 4. Prior Knowledge type of each Web Site.

The test regarding extracting web documents by using the prior knowledge was performed for the web sites in table 1. As we can see from Fig. 5, the most popular type among the types of prior knowledge for the connected web document was serial number type, and the next most popular types was mixed types like serial numbers and image. Also the success rate for using prior knowledge to extract web documents was 100%.

7. Conclusion

This chapter is regarding a conversion system using VoiceXML to provide the HTML contents vocally through a mobile terminal or phone for text disabilities. To decrease the cost and time needed to convert to VoiceXML, the HTML to VoiceXML conversion agent was designed and implemented that is capable of automatically converting HTML documents into VoiceXML documents.

Because it is unable to figure out the meaning of the information on HTML documents, the convertible HTML document type is provided and a hybrid sequential contents selection method was suggested to solve the problem of the stiffness of using structural characteristics and that fail to reflect the user's intention and its effectiveness was proved. Also the contents were separated and extracted to be converted into VoiceXML documents according to the vocal scenario by analyzing the document structure. In addition, the same type of contents across several web documents was extracted at once through the prior knowledge regarding the web linking structure to create practical and effective vocal scenario.

The function and performance of the developed converter were tested on about 400 Korean web pages. In the results, all of the web pages that applied HTML grammar correctly were able to be converted into VoiceXML documents. Therefore, it is determined the application is effective, practical and success rate.

A few web pages that didn't follow exact HTML grammar, especially, if there was no ending tag for the starting tag, were not converted correctly, and to solve this problem, the HTML should be converted into XHTML documents that kept the XML format, and then converted into VoiceXML documents. The presently implemented converter is limited to the convertible HTML documents type because of the inherent problems, so a more intelligent method should be combined to expand the subjects of the conversion.

8. References

- Asakawa, et al, (1998). User Interface of a Homepage Reader, Pro. of ASSET'98, pp149-156.
- Asakawa (2000). Annotation-Based Transcoding for Nonvisual Web Access, Pro. of ASSETS'00, pp172-179
- Buttler, D. ; Liu, L. & Pu C. (2001). A Fully Automated Extraction System for the World Wide Web", IEEE ICDCS-21, pp16-19
- Choi, H.I. & Jang, Y.G. (2001). Design and Implementation of a HTMLtoVoiceXML Converter" , Journal of KISS : Computing Practices, Vol. 7, No. 6, pp559-568.
- Embley D.W.; Jiang Y.S. & Ng Y.K. (1999). Record- boundary discovery in Web documents, Proceedings of the 1999 ACM SIGMOD International Conference on Management of Data , pp467-478

- Embley, D. W. et al, (2006). Notes contemporary table recognition", DAS 2006. LNCS 3872, pp. 164-175
- Goble, C. & Bechhofer, S. et al, (2001). Conceptual Open Hypermedia = The Semantic Web?, Proceedings of the 2nd Int. Workshop on the Semantic Web
- Goose, S.; Newman, M.; Schmidt, C. & Hue, L. (2000). Enhancing Web accessibility via the Vox Portal and a Web-hosted dynamic HTML<->VoxML converter, WWW9, Vol. 33, No. 1-6, pp583-592.
- Hori M.; Kondoh G.; Ono K., Hirose S. & Singhal S. (2000). Annotation-based web content Transcoding, Proc. of WWW9, pp. 197-211
- Huang, A. W. (2000). A Semantic Transcoding System to Adapt Web Services for Users with Disabilities", Pro. of ASSETS'00, pp156-163
- Krulwich, B. (1997). Automating the internet: agents as user surrogates, IEEE Internet Computing, Vol. 1, No. 4, pp. 34-38
- Li, S. et al, (2004). Automatic HTML to XML conversion, WAIM 2004, LNCS 3129, pp. 714-719
- Lieberman, H.; Nardi, B. A. & Wright D. (1999). Training Agents to Recognize Text by Example, Proceedings of Agents'99, pp1-5
- Mohan, R.; Smith, J. & Li, C.-S. (1999). Adapting multimedia internet content for universal access, IEEE Transactions on Multimedia, Vol. 1, No. 1, pp104-114.
- Vankayala, R. R. & Shi, H. (2006). Dynamic voice user interface using VoiceXML and Active Server Pages, APWeb 2006, LNCS 3481, pp. 1181-1184
- W3C (2000). Voice eXtensible Markup Language(Voice XML) version 1.0, <http://www.w3.org/TR/voicexml>, W3C Note 05
- Freire J.; Kumar B. & Lieuwen D. (2001). WebViews: Accessing Personalized Web Content and Services, WWW10, pp576-586
- Luger, G.F. & Stubblefield, W.A. (1997). Artificial Intelligence: Structures and Strategies for Complex Problem Solving, Third Edition. Addison Wesley Longman, Inc.



Assistive Technologies

Edited by Dr. Fernando Auat Cheein

ISBN 978-953-51-0348-6

Hard cover, 234 pages

Publisher InTech

Published online 16, March, 2012

Published in print edition March, 2012

This book offers the reader new achievements within the Assistive Technology field made by worldwide experts, covering aspects such as assistive technology focused on teaching and education, mobility, communication and social interactivity, among others. Each chapter included in this book covers one particular aspect of Assistive Technology that invites the reader to know the recent advances made in order to bridge the gap in accessible technology for disabled or impaired individuals.

How to reference

In order to correctly reference this scholarly work, feel free to copy and paste the following:

Young Gun Jang (2012). An Implementation of Intelligent HTMLtoVoiceXML Conversion Agent for Text Disabilities, Assistive Technologies, Dr. Fernando Auat Cheein (Ed.), ISBN: 978-953-51-0348-6, InTech, Available from: <http://www.intechopen.com/books/assistive-technologies/an-implementation-of-intelligent-htmltovoicexml-conversion-agent-for-text-disabilities>

INTECH
open science | open minds

InTech Europe

University Campus STeP Ri
Slavka Krautzeka 83/A
51000 Rijeka, Croatia
Phone: +385 (51) 770 447
Fax: +385 (51) 686 166
www.intechopen.com

InTech China

Unit 405, Office Block, Hotel Equatorial Shanghai
No.65, Yan An Road (West), Shanghai, 200040, China
中国上海市延安西路65号上海国际贵都大饭店办公楼405单元
Phone: +86-21-62489820
Fax: +86-21-62489821

© 2012 The Author(s). Licensee IntechOpen. This is an open access article distributed under the terms of the [Creative Commons Attribution 3.0 License](https://creativecommons.org/licenses/by/3.0/), which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

IntechOpen

IntechOpen