

異質數位學習內容整合

曾世邦

tsp@tajen.edu.tw

資訊工程與娛樂科技系,
大仁科技大學

Outline

1 Introduction

2 Related Works

- Learning theories in e-learning
- Personalization of learning

3 Proposed method

- Integration by using Vector Space Model
- Improved Integration by Weighted Content

4 Experimental Results

5 Applications

- Test Sheet Assembling

6 Conclusion

E-learning

- **Learning** is the major human activity to accommodate the environment or the society better.
- The learning efficiency becomes more important due to the rapid changing of modern society.
- E-learning is proposed and developed to enhance the human learning efficiency.

Characteristics of the e-learning

- The interaction between teachers and learners are enhanced by the Internet technology.
- The learning activities can be independent from the constraints of time and space.
- There is the higher individuality of learning to the learner.
- The teachers and learners can receive more and faster feedbacks of the learning activities.

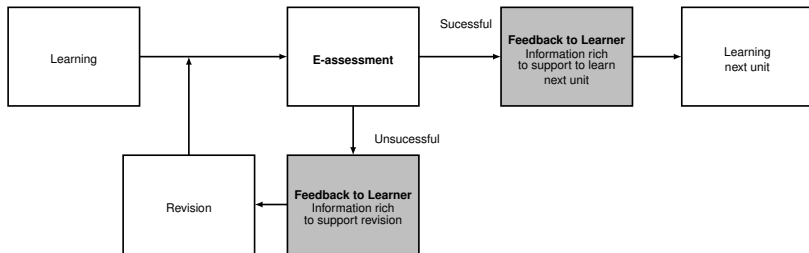
Importance of assessment

- For the personalization, it is necessary to gather information about the learners by assessment process.

assessment

- **Learner:** change the learning behavior.
- **Teacher:** understand the effectiveness of the learning activities to different learners.
- **School:** understand the performance of the educational program.

Relationship between e-assessment and feedback



E-assessment

paper-pencil test (PPT)

- Higher cost
- Less performance

Computer-based test (CBT), e-assessment

- Lower cost
- Higher performance
- Greener

Effectiveness of E-assessment

Test sheet assembling

- selecting the candidate items from the itembank to generate the test sheet

Quality of itembank

- well-organized structure
- **content-richness**

heterogeneous itembanks

One course/textbook taught/edited by many instructors/editors

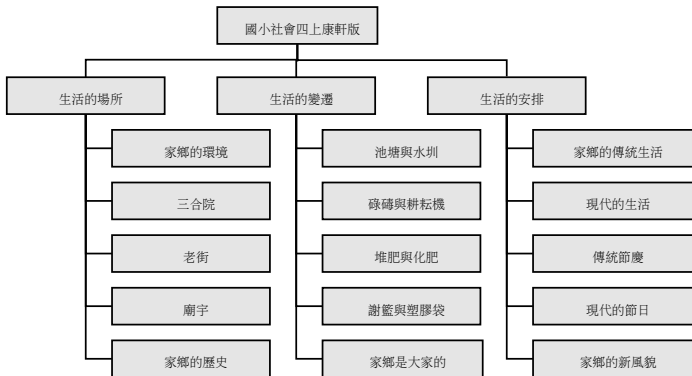
- The instructors/editors of the same course may have different background

One course using several reference books

- The instructors base their teaching material on several reference books
- Learners find related information from other books or articles.

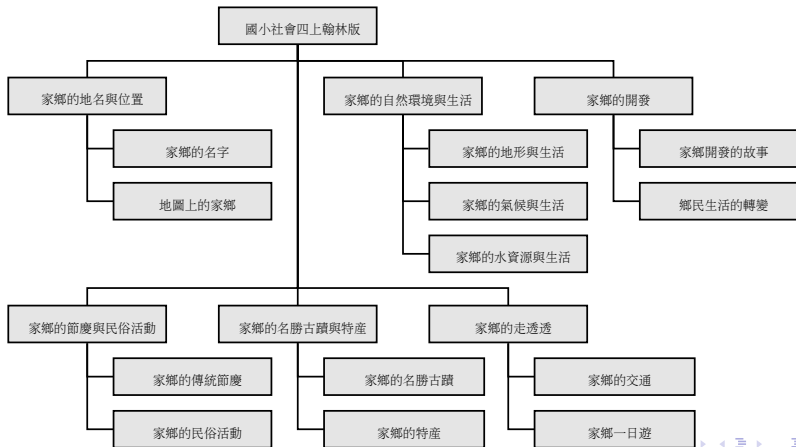
Example: Heterogeneous Itembanks

國小四上社會康軒版



Example: Heterogeneous Itembanks

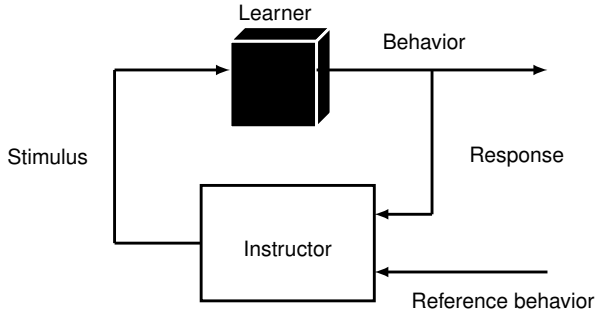
國小四上社會翰林版



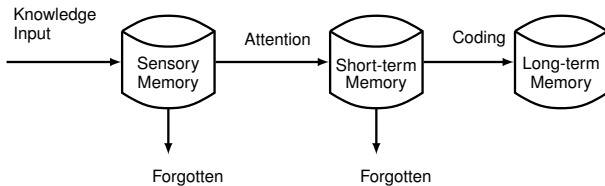
Learning theories in e-learning

- Behaviorism
- Cognitivism
- Constructivism

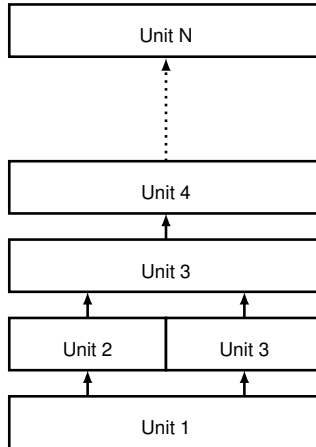
Behaviorism



Cognitivism

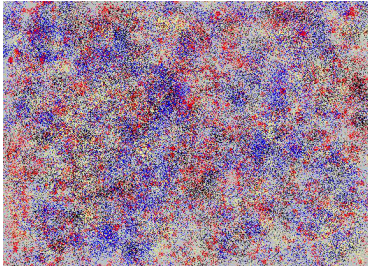


Cognitivism

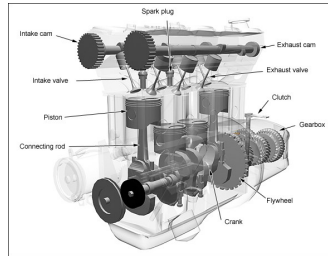


Which unit is the first?

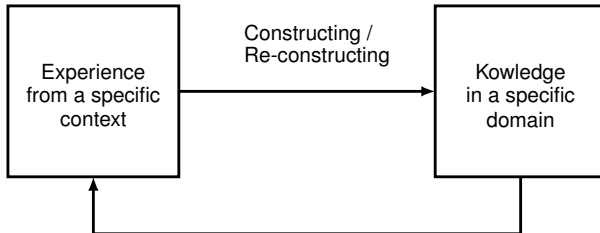
Entropy



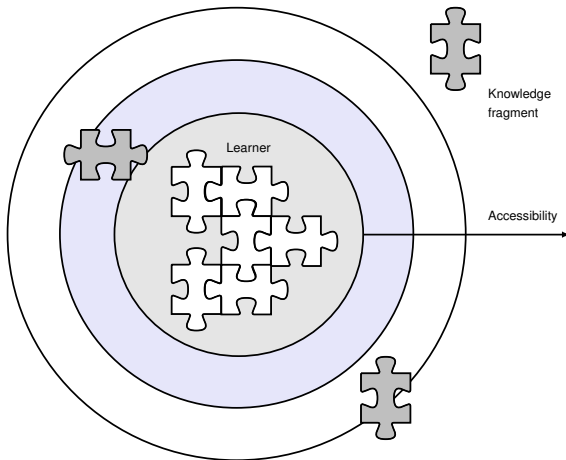
Engine



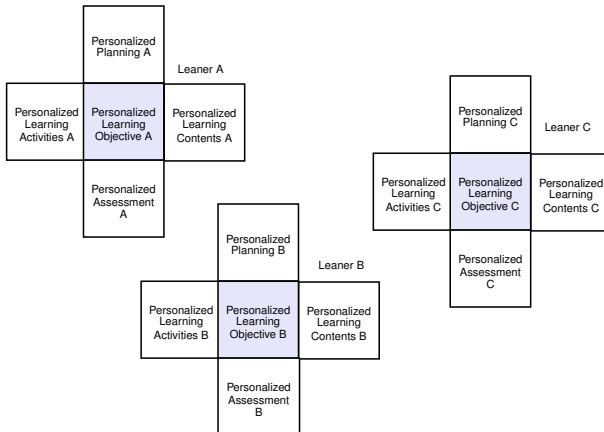
Constructivism



Constructivism



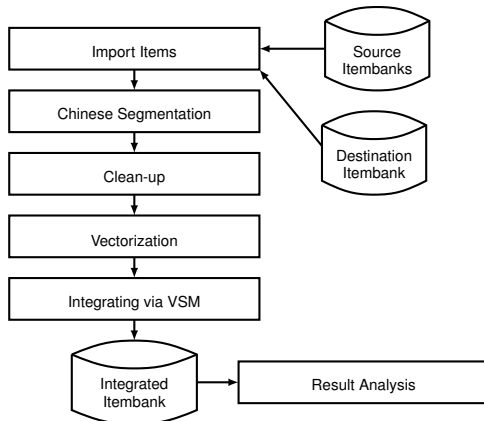
Personalization of learning



Integration of Heterogeneous itembanks

- Integration by using Vector Space Model
- Improved Integration by Weighted Content

Heterogeneous Itembanks Integrator (HIBI)



Data Cleanup

- Chinese segmentation
 - CKIP Chinese Word Segmentation System
<http://ckipsvr.iis.sinica.edu.tw/>
 - Accuracy: 95%-96%
- Removing the stopwords
 - Noun
 - Verb
 - Adjective

Example of Data Cleanup

- 1 item:
通常人們會稱自己()爲自己的家鄉[居住的縣市鄉鎮][工作的地點][結婚地點][求學地]
- 2 segmentation:
通常(D) 人們(Na) 會(D) 稱(VG) 自己(Nh) 爲(P) 自己(Nh) 的(DE) 家鄉(Nc) 居住(VA) 的(DE) 縣市(Nc) 鄉鎮(Nc) 工作(Na) 的(DE) 地點(Na) 結婚 (VA) 地點(Na) 求學(VA) 地(DE)
- 3 extraction:
人們(Na) 稱(VG) 自己(Nh) 自己(Nh) 家鄉(Nc) 居住(VA) 縣市(Nc) 鄉鎮(Nc) 工作(Na) 地點(Na) 結婚(VA) 地點(Na) 求學(VA)

Vector space model

- Documents are represented as vectors.

$$d_j = (w_{1,j}, w_{2,j}, \dots, w_{t,j}) \quad (1)$$

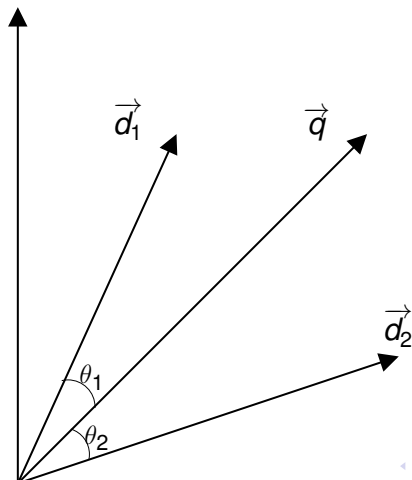
- Each dimension corresponds to a separate term.
- The term specific weights in the document vectors are products of *term frequency* (TF) and *inverse document frequency* (IDF).

$$w_{t,d} = tf_{t,d} \cdot \log \frac{|D|}{|t \in D|} \quad (2)$$

- Document *similarity* is the cosine of the angle between the vectors.

$$similarity = \cos \theta = \frac{\vec{d}_1 \cdot \vec{q}}{|\vec{d}_1| |\vec{q}|} \quad (3)$$

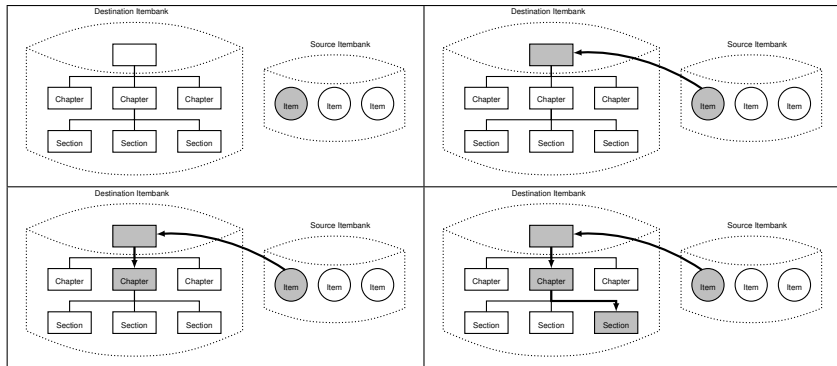
Vector space model



Integrating via VSM

```
1: for each Item  $t$  in the source itembank do
2:    $Depth = 0$ 
3:    $Current = \text{root}$  of the destination itembank
4:   while Node  $Current$  is not a leaf-node do
5:     for each child  $c$  of the  $Current$  do
6:       evaluate the similarity  $Sim(c, t)$ 
7:       if  $Sim(c, t) > MaxSim$  then
8:          $MaxSim = Sim(c, t)$ 
9:          $MaxSimNode = c$ 
10:      end if
11:    end for
12:     $Current = MaxSimNode$ 
13:     $Depth = Depth + 1$ 
14:  end while
15:  Insert item  $t$  into the node  $Current$ 
16: end for
17: Output the destination itembank as integrated itembank
```

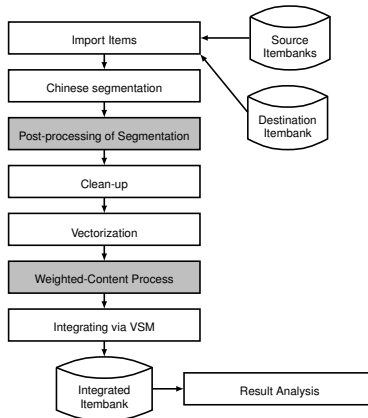
Integrating via VSM



Improved Integration by Weighted Content

- Weighted Content
 - based on the type and structure of the items
 - different parts of the item with different weight
- Post-process of Chinese segmentation

Improved Integration by Weighted Content



Chinese segmentation post-processing

Example of compound noun

” 人口分布 ” $\Rightarrow$ ” 人口 ” , ” 分布 ”

- **Rule 1:** if one adjective directly followed by one noun, the adjective and the noun would be composed into a compound noun.
- **Rule 2:** if one noun directly followed by one noun, the two nouns would be composed into a compound noun.

True/false items

- **True-type:** all the item are with normal weights because the stem contains no irrelevant terms.
- **False-type:** Because the answer is false, it means that there are some irrelevant or error terms in the stem. The detail of answer would be involved into the classifying process and with double weight relative to the stem.

Multiple choice items

- **Stem:** double weight.
- **Options:**
 - **Positive-type:**
 - **Key:** since this option is true/correct, the terms in the key are with double weight.
 - **Distractors:** since this option can contain some error or irrelevant terms, the terms in the distractors are with normal weight.
 - **Negative-type:** identified by "錯"(wrong), "非"(false), "不"(not), "誤"(error), "無"(no), "沒"(no), and "否"(not).
 - **Key:** since this option can contain some really error or irrelevant terms, the terms in the key are with normal weight.
 - **Distractors:** since these options are really true/correct, the terms in the key option are with double weight.

Itembanks

	Society		Nature
	1st semester	2nd semester	2nd semester
Nani	NA	1529	270
Han Lin	697	950	699
Kang Hsuan	464	1106	461

Itembank manager of HGLS



Itembank manager of HGLS

試題編輯 - Windows Internet Explorer

http://hglis.nyu.edu.tw/hglisitembankedit_choices.php hglis

檔案(F) 編輯(E) 格式(O) 我的最愛(S) 工具(T) 說明(H)

☆ ☆ 試題編輯

HGLS 高速網路學習系統
Hyper Global Learning System

測試版

公告事項 題目倉庫 試卷編輯 課程資訊 討論區 網路硬體 查出 (前)

題目倉庫

新增試題 匯出試題 匯入試題 返回

四小 四上 社會 科特 家鄉的地名與位置 家鄉的名字 單元代碼: 629 選擇題共計: 13 題 查詢結果: 13 題

查詢

分頁瀏覽模式

編者		題目內容
編者	18212	通常人們會稱自己()為自己的家鄉。【居住的城市鄉鎮】【旅遊地點】【求學地】
編者	18213	人們通常會在()的地方,形成聚落。【交通不便】【土地貧瘠】【取水方便】【地勢最高】
編者	18214	臺中縣、東勢鎮的地名由來是「因為位於臺原東方的地區,所以取名為東勢角」,由此可知東勢地名的命名原則是依據【地理位置】【氣候】【地形】【人文因素】
編者	18215	下列何者不算聚落?()青山部落 【臺北市】【淡水河口】【臺中縣 梧棲鎮】【茂林鄉】
編者	18216	家鄉的命名可能會依據?()來命名【當地居民的建議】【神明指示】【鄉民的名字】【先人捉鬼的事蹟】
編者	18217	新竹縣、五峰鄉的命名是依據【自然特徵】【人文特徵】【地名的演變】【以上皆非】
編者	18218	臺灣縣、官田鄉的命名是依據【自然特徵】【人文特徵】【地名的演變】【以上皆非】
編者	18219	下列哪一項不是容易形成聚落的條件?【土地肥沃】【水源充足】【交通方便】【地勢最高】
編者	18220	家鄉的命名會根據某些特徵,下列何者不是?【地形】【氣候】【求神問卜】【動、植物的分布】
編者	18221	依照中央氣象局的統計,恆春的月均溫都在()度以上,沒有明顯的季節變化。()應該填入的數字是【0】【10】【20】【30】
編者	18222	小魯的家鄉在花東縱谷平原上,請問他的家鄉可能是什麼型態的聚落?【工業區】【都會區】【農村】【漁村】
編者	18223	屏東市因為位於高雄市、半屏山的東邊而得名,下列哪個位置分布是正確的? ☐ ☐ ☐ ☐

網路網站 100%

Distinguishability

Measure of integration performance

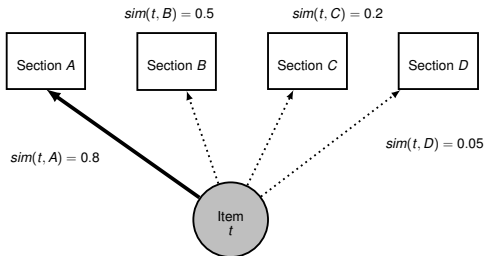
$$\delta = \frac{|\text{sim}(t, a) - \text{sim}(t, b)|}{\max(\text{sim}(t, a), \text{sim}(t, b))} \quad (4)$$

where the two nodes a and b in the meta-itembank are most similar to the item t . If $\delta < 0.2$, the item τ is assumed to be indistinguishable. Otherwise, it is distinguishable.

Distinguishable Ratio

$$DR = \frac{N_d}{N_d + N_i} \times 100\% \quad (5)$$

A example of distinguishable item



$$\delta = \frac{|0.8 - 0.5|}{\max(0.8, 0.5)} = \frac{0.3}{0.8} = 0.375 > 0.2$$

The Distinguishable ratio of integrating different heterogeneous itembanks

Semester	Course	Source	Destination	HIBI _{flat}	HIBI _{tree}	IWC _{flat}	IWC _{tree}
1	Society	Kang Hsuan	Han Lin	79.10	93.53	96.12	96.34
1	Society	Han Lin	Kang Hsuan	78.80	87.37	91.39	90.96
2	Society	Kang Hsuan	Han Lin	72.60	89.33	92.50	93.22
2	Society	Nani	Han Lin	75.93	89.99	92.61	92.94
2	Society	Han Lin	Kang Hsuan	72.21	85.16	89.89	89.89
2	Society	Nani	Kang Hsuan	72.99	83.52	86.52	86.07
2	Society	Han Lin	Nani	71.52	89.96	91.23	91.23
2	Society	Kang Hsuan	Nani	76.11	88.74	92.11	91.16
2	Nature	Kang Hsuan	Han Lin	76.36	89.37	90.89	90.67
2	Nature	Nani	Han Lin	77.78	90.00	92.59	92.59
2	Nature	Han Lin	Kang Hsuan	71.53	82.55	85.98	85.98
2	Nature	Nani	Kang Hsuan	74.81	83.33	86.30	84.07
2	Nature	Han Lin	Nani	78.97	83.69	86.27	87.70
2	Nature	Kang Hsuan	Nani	78.09	84.82	89.80	89.15
Average				75.49	87.24	90.30	90.14
Standard deviation				2.86	3.37	3.05	3.31

Accuracy

Measure of integration performance

- The evaluation of accuracy only takes the distinguishable items into account.
- For the integrating different itembanks, it must be done by the professional expert to verify the integrating result of each items. In the experiments, the destination and source itembank is the same. Therefore it is easy to verify whether the item is inserted into the correct (original) leaf-node.

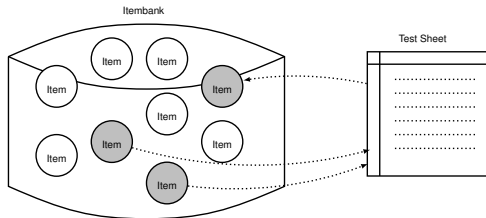
Accuracy

$$AC = \frac{N_{dc}}{N_d} \times 100\% \quad (6)$$

Accuracy of integration

Semester	Course	Source/Destination	HIBI _{flat}	HIBI _{tree}	IWC _{flat}	IWC _{tree}
1	Society	Han Lin	97.40	98.09	99.14	99.73
1	Society	Kang Hsuan	97.60	97.33	97.60	97.62
2	Society	Han Lin	97.39	94.46	97.55	98.08
2	Society	Kang Hsuan	94.99	94.71	96.20	96.48
2	Society	Nani	96.06	94.91	96.13	96.60
2	Nature	Han Lin	98.48	99.57	99.71	99.71
2	Nature	Kang Hsuan	94.06	98.19	97.07	97.29
2	Nature	Nani	99.60	99.99	99.25	99.25
Average			96.95	97.16	97.83	98.09
Standard deviation			1.82	2.21	1.39	1.33

Test Sheet Assembling



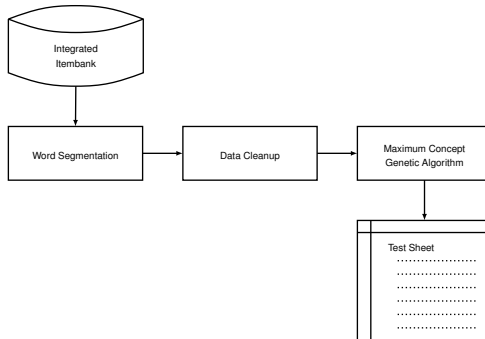
Test Sheet Assembling for Integrated Itembank

Redundancy of itembank

If the itembank aggregates a large amount of items from different sources, the itembank will be with high redundancy.

- The intelligent test sheet assembling is necessary to avoid the redundant items for the integrated itembank.

System architecture of test sheet assembling



Maximum Concepts Genetic Algorithm (MCGA)

- 1: Randomly initiate the population
- 2: **while** The terminate condition is not met **do**
- 3: Selection()
- 4: Crossover()
- 5: Mutation()
- 6: **end while**
- 7: Output the result

The example of crossover in MCGA

	C1					C2				
P1	1	1	0	0	0	1	1	0	1	0
P2	1	0	0	1	1	0	1	0	1	0
O1	1	1	0	1	1	0	1	0	1	0
O2	1	0	0	0	0	1	1	0	1	0

Mutation

Types of mutation

- increment mutation
- decrement mutation

Flip-bit choosing strategy

- Random choosing
- Heuristic choosing
 - increment mutation : increase the maximum fitness value
 - decrement mutation : decrease the minimum fitness value

The example of mutation in MCGA

O1	1	1	0	1	1	0	1	0	1	0
	1	0	0	1	1	0	1	0	1	0
O2	1	0	0	0	0	1	1	0	1	0
	1	0	0	1	0	1	1	0	1	0

Itembanks

	Item#	Keyword#	Identical keywords
DS1	697	7268	1874
DS2	464	3908	1255
DS3	340	3395	1058
DS4	123	1230	443
DS5	357	3873	1144
DS6	157	1707	611
DS7	551	5137	1481
DS8	182	1702	616
DS9	516	5248	1525
DS10	254	2555	887

Coverage

We assume that the concepts in one item can be represented by its keywords. And all identical keywords of the itembank can be as the domain of the itembank. The coverage means that the ratio of the itembank's domain is covered by the test sheet.

$$Coverage = \frac{\text{The identical keywords in the test sheet}}{\text{The identical keywords in the itembank}} \quad (7)$$

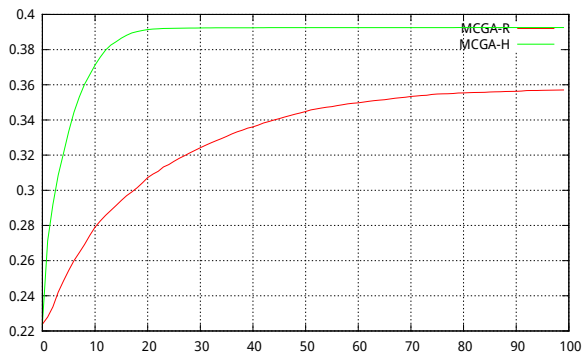
Experimental result of MCGA, size = 50

	Random	MCGA _R		MCGA _H	
	Coverage	Coverage	Time	Coverage	Time
DS1	0.19	0.36	0.78	0.39	5.81
DS2	0.23	0.37	0.50	0.39	1.53
DS3	0.30	0.54	0.42	0.57	1.16
DS4	0.58	0.86	0.16	0.88	0.24
DS5	0.31	0.51	0.46	0.54	1.54
DS6	0.50	0.76	0.23	0.78	0.38
DS7	0.22	0.39	0.62	0.42	3.77
DS8	0.44	0.70	0.23	0.72	0.36
DS9	0.23	0.38	0.63	0.41	3.30
DS10	0.34	0.56	0.34	0.58	0.86

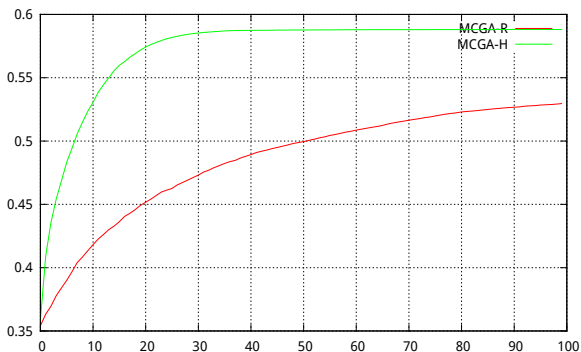
Experimental result of MCGA, size = 100

	Random	MCGA _R		MCGA _H	
	Coverage	Coverage	Time	Coverage	Time
DS1	0.33	0.54	1.56	0.59	10.02
DS2	0.38	0.58	0.97	0.61	3.23
DS3	0.50	0.76	0.80	0.80	1.83
DS4	0.89	1.00	0.30	1.00	0.45
DS5	0.50	0.74	0.87	0.78	2.19
DS6	0.77	0.96	0.42	0.97	0.67
DS7	0.36	0.58	1.20	0.63	5.46
DS8	0.70	0.93	0.43	0.94	0.46
DS9	0.37	0.58	1.24	0.62	5.48
DS10	0.57	0.80	0.65	0.82	1.43

Converging process of MCGA, size = 50



Converging process of MCGA, size = 100



Conclusion

- This thesis focuses on how to integrate the heterogeneous itembanks into an integrated itembank with higher content richness for the e-assessment in the e-learning environment.
- The integrating method is based on the vector space model and improved by the weighted-content strategy and latent semantic indexing.
- The distinguishable ratio is from 75% to 90% in the integration of heterogeneous itembanks.
- In the accuracy tests, the accuracies of distinguishable items are very high, about 97% to 98% on average.

Conclusion

- For the redundancy of the integrated itembank, we developed an effective GA-based test sheet assembling method, called MCGA.
- For more information richness of the assessment feedback, we extended the itembank integrator to the content/item integrator with the ability of integrating the content from Internet, such as wikis and blogs.
- The works are restricted in the descriptive courses, like "Society" and "Nature".

Future work

- A novel and independent schema should be automatic generated to integrate the contents of the collected itembanks
- The sequence in the textbook schema should be considered for the different applications.
- We try to use the multi-objective optimization to assemble the test sheet for more different test requirements.