

These questions do not capture the whole course content!! No guarantees for the answers. Just have fun thinking about these questions. Our exam 2021 was waaaaay easier than SOME of these questions (They asked us to calculate the precision, recall, support confidence. They asked what I asked in Q2 and Q11). I created this file because I was not sure how the exam would look like. Maybe it will be different for your year. Good luck!

Q1:

1. { bread, yogurt, coke }
2. { milk, bread, carrots, coke }
3. { bread, carrots, cheese }
4. { bread, milk }
5. { milk, chocolate, carrots, cheese }
6. { milk, chocolate, yogurt, bread }

Calculate the following:

$S(\text{chocolate} \rightarrow \text{milk})$

$S(\text{coke} \rightarrow \text{bread})$

$c(\text{cheese} \rightarrow \text{carrots})$

And define what the support and confidence are.

Q2:

Give 3 examples of what non-personalised recommender systems could use for recommending stuff (since they can't use content of items, similar users, etc).

Q3:

What are these formulars used for:

- Dice Coefficient
- Pearson correlation
- Cosine Similarity

Q4:

Answer:

What is content-based RS, what is CF RS, and what is knowledge-based RS?

Q5:

This question is about Slope-One predictors.

	Item 1	Item 2	Item 3
Alice	5	?	2
User 1	3	2	3
User 2	-	4	2

Calculate $\text{pred}(\text{Alice}, \text{Item2})$.

Q6:

This question is about probabilistic methods in CB.

Doc-ID	Ayse	Wants	Fries	Label
1	1	1	1	1
2	0	1	1	0
3	1	0	1	0
4	1	1	0	1
5	0	0	1	1
6	1	1	0	?

Calculate $P(X|\text{Label}=1)$.

Q7:

What is the difference between internal validity and external validity in evaluations?

Q8:

Name me the evaluation metrics for classification (excluding the AOC, Precision@K and Recall@K) with their functionality and formular.

Q9:

Imagine you did some predictions. These are the results:

Number of observation	Predicted value	Actual value
1	69	70
2	31	33
3	40	45
4	55	51

Calculate the RMSE.

Q10:

Name me the 3 evaluation settings and the 4 validation strategies mentioned in the lecture.

Q11:

Name me the purposes of explanations with their definition.

Q12:

Name me the Gricean maxims.

Q13:

Name me the 3 stages of NLG Systems.

Q14:

What are conversational interfaces and what challenges could they bring along?

Q15:

When evaluating explanations, we can have 2 different types of evaluations: Extrinsic and Intrinsic. What's their difference?

Q16:

What are the three hybridization designs and what 'sub-types' do they have? Give a short description of each hybridization design and what their sub-types do.

Q17:

Name me all aggregation strategies and their sub-types.

	A	B	C	D	E	F
Ayşe	9	7	3	5	1	10
Zein	6	7	2	3	8	4
Yagmur	10	8	3	2	2	8

Construct the preference list based on:

- Approval voting (Threshold=5)
- Additive
- Least Misery

Q18:

What can be other factors that could influence the groups satisfaction?

Q19:

What are reassuring and repairing explanations and where are they used in Recommender Systems?

Q20:

In context-based RS, what are:

- Explicit Context
- Latent Context
- Static approach
- Dynamic approach
- Complete data
- Incomplete data
- Contextual pre-filtering
- Contextual post-filtering
- Contextual modelling

Define all with 1 sentence. Then name me which are the dimensions of context and which are algorithmic paradigms.

Q21:

What are the common steps when studying fairness? You can just list them.

Q22:

What is the difference between ...

- ... Distributional harms and Representational Harms?
- ... Anti-Classification and Anti-Subordination?

Q23:

Shortly define me:

- Disparate Treatment
- Disparate Impact
- Disparate Mistreatment

Q24:

What is the difference between consumer and provider?

What is the difference between individual and group fairness?

What about other stakeholders?

Q25:

What are reciprocal RS?

Q26:

What are the 4 types of measures when recommending to connect with people?

Q27:

Name me the centrality measures by name only.

Q28:

What is homophily?

Q29:

What are the two diversity metrics and how can they be calculated? What are the meanings of the formulas?

ANSWERS

Q1:

$$S(\text{chocolate} \rightarrow \text{milk}) = 2/6$$

$$S(\text{coke} \rightarrow \text{bread}) = 2/6$$

$$c(\text{cheese} \rightarrow \text{carrots}) = 2/2$$

The support represents the popularity of that product of all products.

Confidence represents the likelihood of purchasing both X and Y.

Q2:

1. Popularity 2. Recency 3. Price

Q3:

- Dice Coefficient: Used to find similarity between items that are represented by keywords
- Pearson correlation: Used to find similarity between users
- Cosine Similarity: Used to find similarity between items

Q4:

Very vague

CB: Uses information of item content and features and users past liked items.

User Profile + Product Features

CF: Uses user-item ratings as input. Assumes that users who had similar taste before will have similar taste in the future. (There is user-user CF and item-item CF)

User Profile + Community Data

KB: Uses knowledge, I guess? There are constraint based and case based models.

User Profile + Knowledge Models

Q5:

$$\text{Pred}(\text{Alice}, \text{item2}) =$$

$$\text{Item3/Item2: } 2-3=-1; 4-2=2 \rightarrow \text{avg. } 0.5 \rightarrow 2+0.5=2.5$$

$$\text{Item2/Item1: } 2-3=-1 \rightarrow 5-1=4$$

$$(2*2.5 + 1*4)/2+1 = \underline{3}$$

Q6:

$$2/3 * 2/3 * 1/3 = 0.149$$

Q7:

Internal: The observed effects are due to the controlled test conditions

External. Results are generalizable

Q8:

Precision for exactness: $TP/(TP+FP)$

Recall for completeness: $TP/(TP+FN)$

Hit Rate to see if at least one item was recommended from y_{test} : 1 if $TP > 1$, and 0 otherwise

F1 is the harmonic balance between precision and recall: $2 * (precision * recall) / (precision + recall)$

Q9:

RMSE: $\sqrt{46/4}$

Q10:

Settings: Lab Studies, Field studies, Offline studies

Validation strategies: Hold-Out, K-Fold, Leave-One-Out, Bootstrap

Q11:

TETPSE:

- Transparency: explain how the system works
- Effectiveness: help making good decisions
- Trust: increase users' confidence
- Persuasiveness: convince to buy or try
- Satisfaction: increase users' satisfaction
- Efficiency: help making decisions faster

Q12:

QQRm:

- Quantity: Be as informative as required, content length
- Quality: Be truthful
- Relevance: Stick to the topic, relevant information
- Manner: Be clear

Q13:

1. Document planning (what to say, content)

2. Microplanning (how to say)

3. Realisation (produce text, say)

Q14:

They support conversation with the users through text, voice or a combination of both. Challenges are high expectations, accents, grammar, low discoverability due to lack of UI

Q15:

Extrinsic: Measure whether system completes the task (helping users in decision making)

Intrinsic: Ask user directly (ratings, feedback)

Q16:

1. Monolithic: aspects of several RS in 1 algorithm

- Feature combination: Combine features from different RS paradigms
- Feature Augmentation: Output of one RS is used to augment the feature space of the actual RS

2. Parallel: two separate RS outputs combined info in the final set

- Mixed: Top scoring next to each other, preference rules
- Weighted: Weighted sum of their scores (the RS should assign ratings on the same scale tho)
- Switching: Switch the RS based on specific situation

3. Pipelined: output of one RS is input of another

- Cascade: based on sequenced order of techniques, successor RS is restricted to predecessor's recommendation list
- Meta level: One RS builds a model that is exploited by the actual RS

Q17:

Majority based: Plural, approval

Consensus based: Additive/avg, multiplicative, fairness

Borderline: Dictatorship, least misery, most pleasure

Results of example:

Approval voting: (A,B),F,(D,E)

Additive: A, (B,F), E, D, C

Least Misery: B, A, F, (C, D), E

Q18:

Emotional contagion, conformity, group attributes (user roles, assertive users, expertise users, dominance)

Q19:

Reassuring ones are used in the case of agreements. Repairing ones are used in the case of disagreement. These types are used in group explanations.

Q20:

- Explicit Context: specified directly (through questionnaires, etc.)
- Latent Context: hidden/inferred
- Static approach: predefined attributes, stable, fixed
- Dynamic approach: different activities give rise to different contexts
- Complete data: All contextual factors are known explicitly (static, explicit)
- Incomplete data: Only a subset of factors is known
- Contextual pre-filtering: Use context to filter out relevant data, then predict
- Contextual post-filtering: Predict first, then filter out using context
- Contextual modelling: Use context while predicting

The last three are the paradigms, the rest are the dimensions of context.

Q21:

Define fairness, collecting data, experimental design, reporting results

Q22:

Distributional harms: If a resource/opportunity is unequally distributed

Representational harms: If people are being represented in an invalid way

Anti-Classification: The protected attribute should not play a role in decision making

Anti-Subordination: Current decision-making process should work to reverse the effect of discrimination

Q23:

Disp. Treatment: sensitive attribute is a direct feature in decision making

Disp. Impact: outcomes are unequally distributed between groups

Disp. Mistreatment: one group has higher FP than another

Q24:

Consumer: The people that receive the recommended list

Provider: The people that provide the items in the recommended list

Other stakeholder: Neither consumer nor provider but they are still affected by the fairness

Individual F: Similar users are treated similarly

Group F: sensitive attributed identify a protected group, discrimination should be avoided

Q25:

In reciprocal RS users have profiles and preferences, success is determined by user, popular item (user) should not be recommended to too many users

Q26:

Content, Content+Link, FoF, SONAR

Q27:

Degree centrality, Betweenness, Closeness, EigenCentrality, PageRank

Q28:

Homophily (similar gender, affiliation, etc.) is the most important factor for predicting collaboration.

Q29:

Long tail novelty: expected popularity complement

- EPC: the expected number of (new) read relevant recommended items

Unexpectedness: expected profile distance

- EPD: an article that is unexpected, but still useful for the reader