



Enhancement of the Neutrality in Recommendation

Toshihiro Kamishima^{*}, Shotaro Akaho^{*}, Hideki Asoh^{*}, and Jun Sakuma[†]

^{}National Institute of Advanced Industrial Science and Technology (AIST), Japan*

[†]University of Tsukuba, Japan; and Japan Science and Technology Agency

Workshop on Human Decision Making in Recommender Systems

In conjunction with the RecSys 2012 @ Dublin, Ireland, Sep. 9, 2012

START

1

Today, we would like to talk about the enhancement of the neutrality in recommendation.

Overview

Decision Making and Neutrality in Recommendation

Decisions based on biased information brings undesirable results



Providing neutral information is important in recommendation



Information Neutral Recommender System

The absolutely neutral recommendation is intrinsically infeasible, because recommendation is always biased in a sense that it is arranged for a specific user,



A system makes recommendation so as to enhance the neutrality from a viewpoint specified by a user

Because decisions based on biased information brings undesirable results, providing neutral information is important in recommendation.
For this purpose, we propose an information neutral recommender system.
Unfortunately, the absolutely neutral recommendation is intrinsically infeasible.
Therefore, this recommender system makes recommendation so as to enhance the neutrality from a viewpoint specified by a user.

Outline

Introduction

Importance of the Neutrality and the Filter Bubble

- ★ influence of biased recommendation, filter bubble problem, discussion in the RecSys 2011 panel

Neutrality in Recommendation

- ★ ugly duckling theorem, information neutral recommendation

Information Neutral Recommender System

- ★ latent factor model, viewpoint variable, neutrality function

Experiments

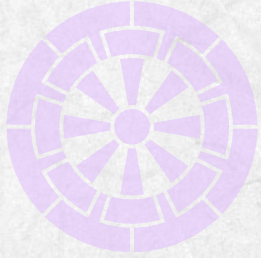
Conclusion

This is an outline of our talk.

After showing the importance of the neutrality in recommendation, we introduce the Filter Bubble problem.

We then discuss the neutrality in recommendation, and show our information neutral recommender system.

Finally, we summarize our experimental results, and conclude our talk.



Importance of the Neutrality and the Filter Bubble



We begin with the importance of the neutrality and the filter bubble problem.

Biased Recommendation

Biased Recommendation

- ✳ exclude a good candidate from a set of options
- ✳ rate relatively inferior options higher



Inappropriate Decision



The Filter Bubble Problem

Pariser posed a concern that personalization technologies narrow and bias the topics of information provided to people

<http://www.thefilterbubble.com/>

Biased recommendations may exclude a good candidate from candidates, or may rate relatively inferior option higher.

Consequently, decisions would become inappropriate.

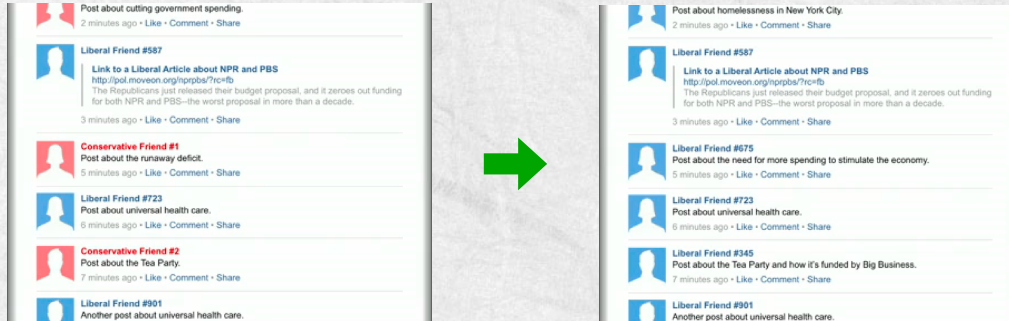
Pariser pointed out a problem of such biased recommendations as the filter bubble problem, which is a concern that personalization technologies narrow and bias the topics of information provided to people.

Filter Bubble

[TED Talk by Eli Pariser]

Friend Recommendation List in Facebook

conservative people are eliminated from Pariser's recommendation list



A summary of Pariser's claim

- ✳ Users lost opportunities to obtain information about a wide variety of topics
- ✳ Each user obtains too personalized information, and this makes it difficult to build consensus in our society

Pariser shows an example of a friend recommendation list in Facebook.

To fit for his preference, conservative people are eliminated from his recommendation list, while this fact is not noticed to him.

His claim would be summarized into these two points.

Users lost opportunities to obtain information about a wide variety of topics.

Each user obtains too personalized information, and this makes it difficult to build consensus in our society.

RecSys 2011 Panel on Filter Bubble

[RecSys 2011 Panel on the Filter Bubble]

RecSys 2011 Panel on Filter Bubble

- ★ Are there “filter bubbles?”
- ★ To what degree is personalized filtering a problem?
- ★ What should we as a community do to address the filter bubble issue?

<http://acmrecsys.wordpress.com/2011/10/25/panel-on-the-filter-bubble/>



Intrinsic trade-off

providing
a diversity of topics



focusing on
users' interests

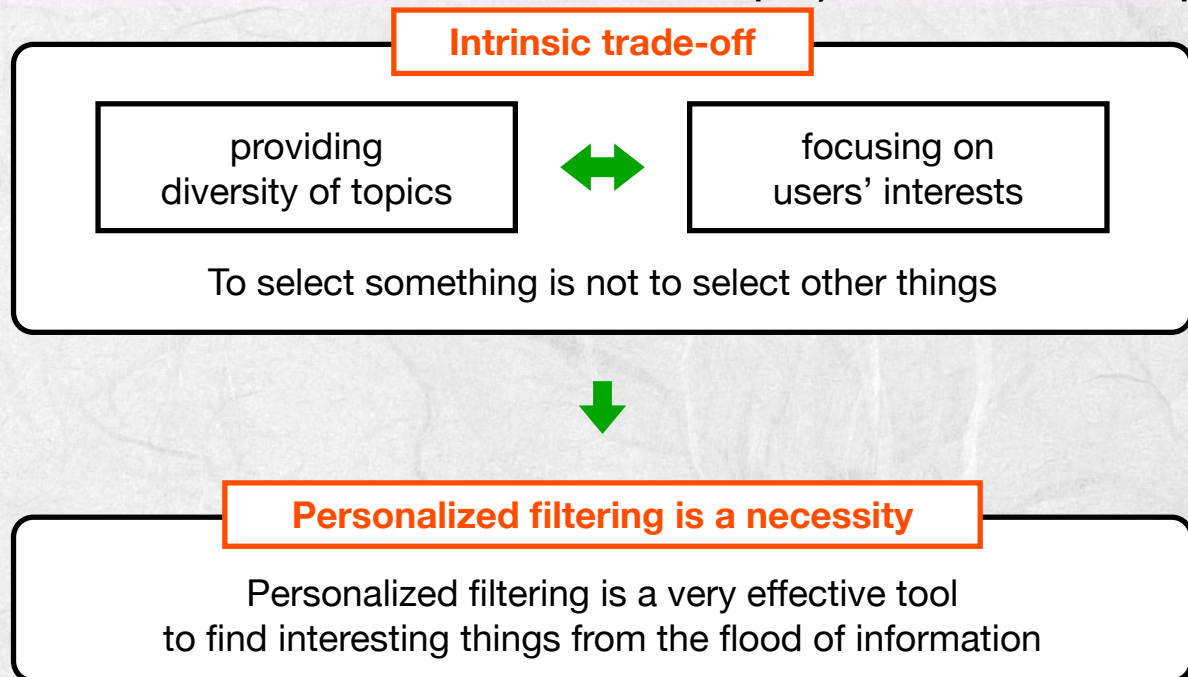
To select something is not to select other things

In the last RecSys 2011 conference, a panel on this filter bubble problem was held. These three sub-problems are discussed.

For the first sub-problem, panelists pointed out that the filter bubble is an intrinsic trade-off between providing a diversity of topics and focusing on users' interests, because to select something is not to select other things.

RecSys 2011 Panel on Filter Bubble

[RecSys 2011 Panel on the Filter Bubble]



8

Though personalized filtering has such a flaw, it is a very effective tool to find interesting things from the flood of information.
Clearly, personalized filtering is a necessity.

RecSys 2011 Panel on Filter Bubble

[RecSys 2011 Panel on the Filter Bubble]

Personalized filtering is a necessity

Personalized filtering is a very effective tool to find interesting things from the flood of information



recipes for alleviating undesirable influence of personalized filtering

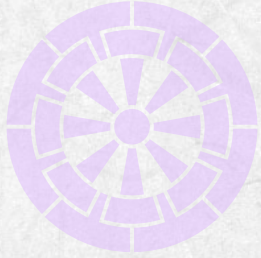
- ★ capture the users' long-term interests
- ★ consider preference of item portfolio, not individual items
- ★ follow the changes of users' preference pattern

our approach

give users to control perspective to see the world through other eyes

In the RecSys panel, panelists suggested recipes for alleviating undesirable influence of personalized filtering.

Among these recipes, we took an approach to give users to control perspective to see the world through other eyes.



Neutrality in Recommendation



We then discuss the neutrality in recommendation.

Ugly Duckling Theorem

[Watanabe 69]

Ugly Duckling Theorem on fundamental property of classification

if the similarity between a pair of ducklings is measured by the number of potential binary classification rules that classifies both of them into the same positive class

similarity between
an ugly and a normal ducklings

=

similarity between
any pair of normal ducklings



An ugly and a normal ducklings are indistinguishable



Extremely unintuitive! Why?

We classify them by methods like SVMs everyday...

11

Before discussing the neutrality, we reconsider the well-known ugly duckling theorem on a fundamental property of classification.

According to this theorem, under this condition, the similarity between a ugly and a normal ducklings is equivalent to the similarity between any pair of normal ducklings.

This fact derives the fact that an ugly and a normal ducklings are indistinguishable.

This looks extremely unintuitive! Why?

Ugly Duckling Theorem

[Watanabe 69]



Extremely unintuitive! Why?



The number of classification rules are considered,
but properties of rules are completely ignored

- ★ **All features are equally treated**

ex. the weight of a *body color* feature equals to that of a length of a duckling

- ★ **The complexity of rules is ignored**

ex. the number of features included in rules



When classification, one must emphasize some features of objects
and must ignore the other features

12

This is because the number of classification rules are considered, but properties of rules are completely ignored: all features are equally treated and the complexity of rules is ignored. This theorem implies that When classification, one must emphasize some features of objects and must ignore the other features.

Information Neutral Recommendation

Ugly Duckling Theorem

A part of aspects must be stressed when classifying objects



**It is infeasible to make recommendation
that is neutral from any viewpoints**

Information Neutral Recommendation

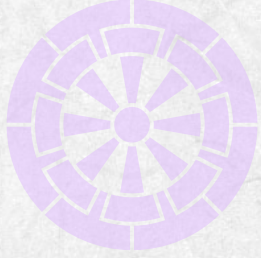
the neutrality from a viewpoint specified by a user
and other viewpoints are not considered

ex. A recommender system enhances the neutrality in terms of whether conservative or progressive, but it is allowed to make biased recommendations in terms of other viewpoints, for example, the birthplace or age of friends

13

Because the ugly duckling theorem indicates that a part of aspects must be stressed when classifying objects, it is infeasible to make recommendation that is neutral from any viewpoints. Therefore, we took an approach of enhancing the neutrality from a viewpoint specified by a user and other viewpoints are not considered.

In the case of Pariser's Facebook example, a system enhances the neutrality in terms of whether conservative or progressive, but it is allowed to make biased recommendations in terms of other viewpoints, for example, the birthplace or age of friends.



Information Neutral Recommender System



To enhance such neutrality, we propose an information neutral recommender system.

Information Neutral Recommender System

v : viewpoint variable

A binary variable representing a viewpoint specified by a user

Information Neutral Recommender System

neutral from a specified viewpoint

maximize statistical independence
between a preference score and a viewpoint variable

+

high prediction accuracy

minimize an empirical error plus a L2 regularization term



Information neutral version of a latent factor model

15

This system adopt a viewpoint variable, which is a binary variable representing a viewpoint specified by a user.

A goal of an information neutral recommender system is to make recommendation that is neutral from a specified viewpoint while keeping high prediction accuracy.

The neutrality is enhanced by statistical dependence between a preference score and a viewpoint variable.

High prediction accuracy is achieved by minimizing an empirical error plus a L2 regularization term.

We then show an information neutral version of a latent factor model.

Latent Factor Model

[Koren 08]

Latent Factor Model : basic model of matrix decomposition

Predicting Ratings Task

predict a preference score of an item y rated by a user x

$$\hat{s}(x, y) = \mu + b_x + c_y + \mathbf{p}_x \mathbf{q}_y^\top$$

Diagram labels for the equation above:

- global bias (points to μ)
- cross effect of users and items (points to $\mathbf{p}_x \mathbf{q}_y^\top$)
- user-dependent bias (points to b_x)
- item-dependent bias (points to c_y)

For a given training data set, model parameters are learned by minimizing the squared loss function with a L_2 regularizer

16

A latent factor model is a basic model of matrix decomposition, and is designed for predicting a preference score.

A preference score is modeled by this formula, which consists of three bias terms and one cross term.

For a given training data set, model parameters are learned by minimizing the squared loss function with a L_2 regularizer.

Information Neutral Latent Factor Model

modifications of a latent factor model

- ★ **adjust scores according to the state of a viewpoint**
incorporate dependency on a viewpoint variable
- ★ **enhance the neutrality of a score from a viewpoint**
add a neutrality function as a constraint term

adjust scores according to the state of a viewpoint

viewpoint variables

$$\hat{s}(x, y, v) = \mu^{(v)} + b_x^{(v)} + c_y^{(v)} + \mathbf{p}_x^{(v)} \mathbf{q}_y^{(v)\top}$$

- ★ Multiple latent factor models are built separately, and each of these models corresponds to the each value of a viewpoint variable
- ★ When predicting scores, a model is selected according to the value of viewpoint variable

17

These two points are modified in information neutral version of a latent factor model.

First, we modify this model so as to be able to adjust scores according to the state of a viewpoint to incorporate dependency on a viewpoint variable.

Multiple latent factor models are built separately, and each of these models corresponds to the each value of a viewpoint variable.

When predicting scores, a model is selected according to the state of viewpoint variable.

Information Neutral Latent Factor Model

enhance the neutrality of a score from a viewpoint

neutrality function, $\text{neutral}(\hat{s}, v)$: quantify the degree of neutrality

The larger output of a neutrality function,
the higher degree of the neutrality of a prediction score
from a viewpoint variable

neutrality parameter to balance
between the neutrality and the accuracy

regularization
parameter

$$\sum_{\mathcal{D}} (s_i - \hat{s}(x_i, y_i, v_i))^2 - \eta \text{neutral}(\hat{s}(x_i, y_i, v_i), v_i) + \lambda \|\Theta\|_2^2$$

squared loss function neutrality function L₂ regularizer

Parameters are learned by minimizing this objective function

18

Second, a model is modified so as to be able to enhance the neutrality between a score and a viewpoint.

For this purpose, we introduce a neutrality function to quantify the degree of neutrality.

This neutrality function is added to the objective function like a regularization term.

A neutrality parameter η balances between the neutrality and the accuracy.

Parameters are learned by minimizing this objective function.

Mutual Information as Neutrality Function

neutrality = scores are not influenced by a viewpoint variable



neutrality function = negative mutual information

We treat the neutrality as the statistical independence,
and it is quantified by mutual information
between a predicted score and a viewpoint variable

$\Pr[\hat{s}|v]$ is required for computing mutual information



This distribution function modeled by a histogram model



Failed to derive an analytical form of gradients of objective function



An objective function is minimized by a Powell method without gradients

19

We finally formalize a neutrality function.

Here, the neutrality means that scores are not influenced by a viewpoint variable.

Therefore, we treat the neutrality as the statistical independence, and it is quantified by mutual information between a predicted score and a viewpoint variable.

The computation of mutual information is fairly complicated, but we here omit the details.



Experiments

20

We finally summarize our experimental results.

Experimental Conditions

General Conditions

- ★ 9,409 use-item pairs are sampled from the Movielens 100k data set (A Powell optimizer is computationally inefficient and cannot be applied to a large data set)
- ★ the number of latent factor $K = 1$
- ★ regularization parameter $\lambda = 0.01$
- ★ Evaluation measures are calculated by using five-fold cross validation

Evaluation Measure

- ★ **MAE (mean absolute error)**
 - ★ prediction accuracy
- ★ **NMI (normalized mutual information)**
 - ★ the neutrality of a preference score from a specified viewpoint (mutual information between the predicted scores and the values of viewpoint variable, and it is normalized into the range $[0, 1]$)

21

These are our experimental conditions.

We tested on this sampled data set, because a Powell optimizer is computationally inefficient and cannot be applied to a large data set.

We used two types of evaluation measure.

MAE, mean absolute error, measures prediction accuracy.

NMI, normalized mutual information, measures the neutrality between a predicted score and a viewpoint variable.

Viewpoint Variables

The values of viewpoint variables are determined depending on a user and/or an item

“Year” viewpoint : a movie’s release year is newer than 1990 or not

The older movies have a tendency to be rated higher, perhaps because only masterpieces have survived [Koren 2009]

“Gender” viewpoint : a user is male or female

The movie rating would depend on the user’s gender

22

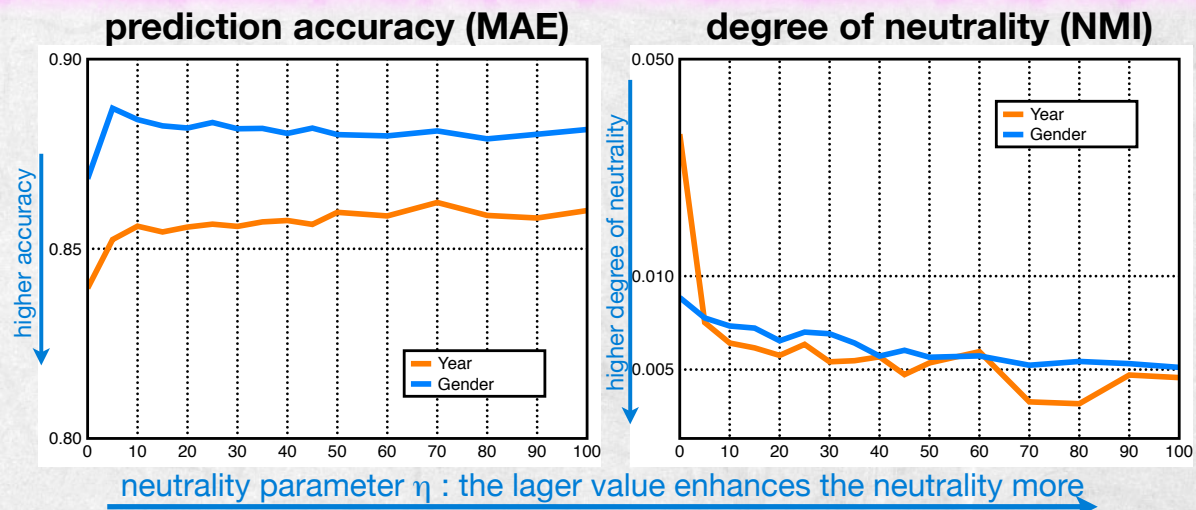
We tested two types of viewpoint variables.

The values of viewpoint variables are determined depending on a user and/or an item.

First, a “Year” viewpoint variable represents whether a movie’s release year is newer than 1990 or not.

Second, a “Gender” viewpoint variable represents a user is male or female.

Experimental Results



As the increase of a neutrality parameter η , prediction accuracies were worsened slightly, but the neutralities were improved drastically

INRS could successfully improved the neutrality without seriously sacrificing the prediction accuracy

These are our experimental results.

X-axes correspond to neutrality parameters, the lager value enhances the neutrality more.

This chart (left) shows the change of prediction accuracy.

This chart (right) shows the change of the degree of neutrality.

As the increase of a neutrality parameter η , prediction accuracy worsened slightly, and the neutrality enhanced drastically.

Therefore, we can conclude that our information neutral recommender system could successfully improved the neutrality without seriously sacrificing the prediction accuracy.

Conclusion

Our Contributions

- ✳ We formulate the neutrality in recommendation based on the ugly duckling theorem
- ✳ We developed a recommender system that can enhance the neutrality in recommendation
- ✳ Our experimental results show that the neutrality is successfully enhanced without seriously sacrificing the prediction accuracy

Future Work

- ✳ Our current formulation is poor in its scalability
 - ✳ formulation of the objective function whose gradients can be derived analytically
 - ✳ neutrality functions other than mutual information, such as the kurtosis used in ICA
- ✳ Information neutral version of generative recommendation models, such as a pLSA / LDA model

24

These are our contributions.

Our current formulation is poor in its scalability.

We plan to develop the formulation of the objective function whose gradients can be derived analytically.

We also consider to use a neutrality function other than mutual information, such as the kurtosis used in ICA.

program codes and data sets

<http://www.kamishima.net/inrs>

acknowledgements

- ✳ We would like to thank for providing a data set for the Grouplens research lab
- ✳ This work is supported by MEXT/JSPS KAKENHI Grant Number 16700157, 21500154, 22500142, 23240043, and 24500194, and JST PRESTO 09152492

Program codes and data sets are available at here.
That's all I have to say. Thank you for your attention.