

대형 쇼핑몰에서 K-means 알고리즘과 순차 패턴을 이용한 추천 시스템

Recommendation System based on Sequential Pattern and K-means Algorithm for Big Shopping Mall

심장섭¹, 정순기²

Jang Sup Sim¹, Sun Ki Jung^{2*}

요약

본 연구에서 제안된 시스템은 매일 변화하는 고객의 관심도를 반영하기 위해, 매일 고객을 재그룹화 하고, 실시간 추천에 필요한 배치 작업을 수행한다. 고객군 그룹화 작업은 K-means 군집화 알고리즘을 사용하고, 각 고객 그룹별로 관심 상품의 시계열적인 변화를 순차패턴 알고리즘을 사용하여 추출해낸 후, 실시간으로 고객별 과거 관심 상품 내역과 배치 작업으로 추출한 순차 패턴을 비교하여 추천 상품을 결정한다. 이를 통해 대용량의 계층 구조를 가지면서 고객의 상품에 대한 관심 정보가 일단위로 연속해서 변화하는 stream data 환경에 대한 새로운 추천 프레임워크를 제시했으며, 대형 쇼핑몰의 실제 데이터를 활용하여 그 효용성을 보였다.

핵심어: 추천시스템, K-means 알고리즘, 데이터 마이닝, 데이터모델, 빅쇼핑몰

Abstract

The recent study trend in data-mining technique is progressing toward the formalization of user profile, association rule, sequential pattern filtering, similar time sequence discovery, and many other hybrid methodologies of various techniques. The data for modeling can use and extract a form logged file, but the user information that has been already available as to age and gender is not utilized. Due to the privacy issues and data reliability of information on the Internet, the system is designed to perform a predictive modeling of user database segmentation by user's interest and preference of the given category information. A user preference defines the user's feature vector that is based on preferred category information, and user

1 Department of Computer Engineering, chung Buk National Univ, chungdae-ro1 seowon-gu chungju-city, Chungbuk-Do, 28644, Korea

e-mail: jssim@chungbuk.ac.kr

2 Department of Computer Engineering, chung Buk National Univ, chungdae-ro1 seowon-gu chungju-city, Chungbuk-Do, 28644, Korea

e-mail: sk-jung@chungbuk.ac.kr (Corresponding author)

Received(May 30.2014), Review (June 16.2014), Accepted(June 30.2014)

database Clustering is done with the conclusive result as the process. This paper requires construction of sequential database in order to extract the pattern of users. In other words, establishing the database by user segment. The methodology of test has been performed to carry out by the first half month of gathering test data and performed to the reliability for recommending from evaluation system used on the remaining second half of the month of that data. The recommendation engine through sequential pattern extracted technique of which suggested by the large capacity hierarchical structured data was able to confirm the accuracy of recommendation and a fast recommendation time of $O(N)$ as well as learning time.

Keyword: recommendation system, K-means algorithm, data mining, data model, Big shopping mall.

1. 서론

오늘날 온라인 쇼핑몰에서 상품은 구매하는 입장에서는 다양하고 많은 상품을 비교해 가면서 상대적으로 저가이면서, 물리적인 시간과 장소를 구애받지 않고 구매할 수 있는 시대가 되었다. 그러나 오프라인과 달리 온라인 쇼핑몰은 고객의 구매를 유도하거나 안내할 판매사원이 존재하지 않기 때문에 고객의 취향에 맞는 적합한 상품을 찾기 위해서는 많은 정보를 처리해야만 한다. 이러한 정보의 과부하(Information Overload)는 고객의 구매의욕을 상실시키는 문제 중 하나가 된다[1].

따라서 본 논문에서는 대용량 계층 구조를 갖는 스트림데이터 환경에서 효과적으로 활용이 가능한 추천 시스템 프레임 워크를 제안하고 검증하는데 목적이 있다. 제안하는 목표 시스템은 대형 쇼핑몰에서 고객의 선호도에 맞게 상품을 추천하는 추천 시스템으로써, 주요 기술적인 특징은 다음과 같다.

첫째, 상품수가 수십만, 고객수가 수백만 이상인 환경에서 동작할 수 있는 구조를 설계한다. 국내 온라인 쇼핑몰의 상황에 맞는 양의 데이터를 처리할 수 있는 추천 프레임 워크를 제시한다. 널리 알려져 있고, 안정적인 실행 속도를 가지고 있는 것으로 입증된 K-means 군집화와 순차패턴 알고리즘을 활용한다.

둘째, 매일 변화하는 고객의 관심도를 반영한다. 대형 인터넷 사이트에서는 실시간 분석이나 처리가 불가능 하기 때문에 매일 batch 처리를 통해 분석을 수행한다. 제안된 시스템은 매일, 매일 주기적으로 쌓이는 각 고객의 카테고리, 상품에 대한 접근 정보를 활용하여, 고객이 다시 방문했을 때, 고객이 선호하는 상품을 추천하기 위해 데이터 마이닝 기법중의 하나인 순차 패턴(Sequential Pattern)을 이용했다.

셋째, 카테고리과 상품간의 관계정보를 반영한다. 대형 인터넷 서비스에서 예외 없이 카테고리

정보를 활용하여 상품에 대한 접근을 카테고리에 대한 관심으로 반영해 준다. 즉, 고객의 상품에 대한 접근 수 및 카테고리에 대한 접근 수를 수치화 하여 카테고리에 대한 관심으로 반영하며, 카테고리에 대한 관심 정보는 고객의 군집화에 활용된다.

2. 관련 연구

2.1 추천 시스템

웹 사이트 개인화란 특정 고객 또는 고객 집단에게 맞춤형 웹 페이지를 제공함으로써 고객의 구매력과 반응성(responsibility)을 높이기 위한 마케팅 전략이다. 웹 사이트 개인화는 고객의 프로파일, 거래 데이터, 웹 로그 데이터 등을 토대로 각 고객의 구매 상황과 특성을 분석하여 각각의 개별 고객에게 적합한 정보를 제공하는 기술로서, 온라인 쇼핑몰에서는 특히 그 중요성이 강조되고 있다[2]. 이러한 개인화 기술 중의 하나인 상품 추천 기술은 고객들에게 추천 상품 리스트를 만들어 고객들이 구매 가능성이 있는 상품을 쉽게 찾도록 도와주는 개인화된 정보 필터링 기술이다. 따라서 상품 추천 시스템은 고객들의 상품구매 가능성을 개선하기 위하여 분석적 기법을 사용하여 적절한 상품을 고객들에게 제공한다. 상품 추천 방법들은 고객에게 개인적인 선호도를 어느 정도까지 명시적으로 표시하도록 요구하느냐에 따라 구분할 수 있다.

가) 협업 필터링(Collaborative Filtering)

협업 필터링은 추천을 수행할 목표 고객과 비슷한 선호도를 보이는 이웃(neighborhood)들이 구매상품 중에서 구매할 가능성이 가장 높은 상품을 추천하는 방법이다. 즉 해당 고객과 선호도가 유사한 고객들의 상품에 대한 평가 점수를 활용하여 고객에게 적합한 상품정보를 추천하는 방법이다. 이 기술은 Goldberg et al.[1992]에 의해서 처음 등장한 개념으로 초기에는 Tapestry라 불리는 e-mail 필터링 도구로 사용되었다. 그러나 이 방법은 사용자의 수동적이고 복잡한 질의를 통해서 접근이 가능한 불편함 때문에 보편화 되지는 않았다. 그 이후 GroupLens는 1994년 neighborhood based 알고리즘을 이용하여 유즈넷 뉴스에 대한 개인화된 예측을 했다. 이것이 현재 널리 쓰이고 있는 협업 필터링 알고리즘의 토대를 닦았다고 할 수 있다. GroupLens 알고리즘은 피어슨(Pearson)의 상관계수(correlation coefficient)를 이용하여 개인화된 예측을 하는 알고

리즘 이다[3][4].

$$P_{a,i} = \bar{r}_a + \frac{\sum_{u=1}^n (r_{u,i} - \bar{r}_u) \times W_{a,u}}{\sum_{u=1}^n W_{a,u}} \quad (1)$$

여기서, $P_{a,i}$ 는 상품 i 에 대한 추천대상 고객 a 의 예측 값이다.

n 은 추천 대상고객 a 에 대응하는 이웃(neighbor)고객의 수이고, $W_{a,u}$ 는 이웃 고객 u 와 추천 대상고객 a 의 유사성(similarity) 가중치(weight)이다.

위의 알고리즘을 이용하여 각 상품에 대한 추천의 대상이 되는 고객의 선호도점수가 계산되고, 이로부터 선호점수가 높은 상품을 고객에게 추천해 준다.

나) 속성기반 필터링(Contents Based Filtering)

속성기반 필터링은 'Information Filtering'이라고도 부르며 고객들이 관심 있어 하는 속성을 파악하여 고객들이 과거에 좋아했던 상품과 가장 유사한 상품을 추천하는 방법이다. 상품과 상품사이의 연관성을 기반으로 추천하는 방식으로 유사한 상품을 찾기 위해서는 각각의 개별 상품의 특성을 명확하게 추출해 내야 하는데 이것이 사람에 의존해야하기 때문에 제대로 이루어지지 않는 한계점이 있다. 즉, 속성 기반 필터링은 상품의 속성을 명확하게 추출하고 관리하기가 어려워 많은 수의 상품을 다루는 대형 온라인 쇼핑몰에서는 현실적으로 적용하기 불가능하다는 단점이 있다[5].

2.2 데이터 마이닝

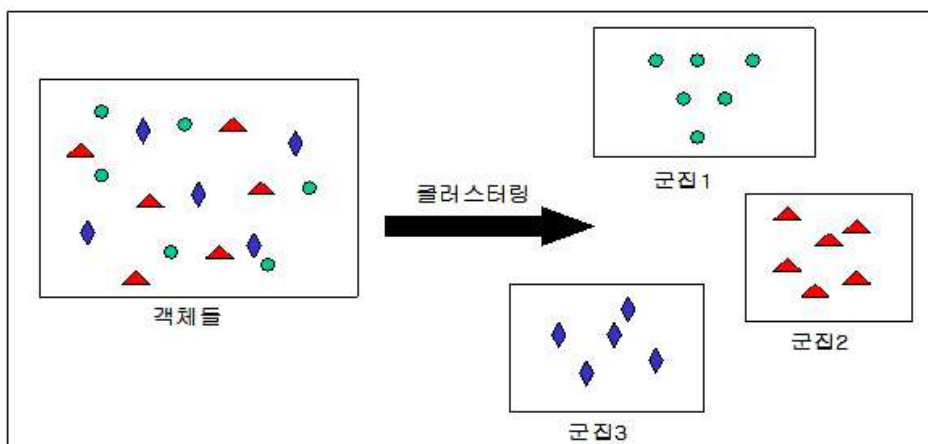
데이터 마이닝 기법은 대용량 데이터에서 쓸모 있는 지식을 찾아내기 위한 기법으로 클러스터링, 연관규칙 등이 있다.

가) 클러스터링(Clustering)

클러스터링(Clustering)은 주어진 데이터 집합을 서로 유사한 군집(cluster 또는 segment)으로 분할해 나가는 과정으로 하나의 군집에 속하는 데이터 간에는 서로 다른 군집내의 데이터와는 구분되는 유사성(similarity)을 갖게 된다. 달리 말하면 클러스터링은 유사한 특징을 갖는 객체들

을 몇 개의 군집으로 생성하는 데이터 마이닝 기법 중의 하나로서 순차패턴 보다는 훨씬 널리 알려지고 보편적인 방법으로 알려져 있다. 때로는 분류작업(classification)과 흔히 혼동 되지만, 군집을 이루는 조건이 사전에 주어지지 않는다는 점에서 분류작업과는 다른 기법이다. 즉, 클러스터링은 분류와 달리 미리 정의된 클래스나 분류규칙을 학습할 예제 데이터가 없다는 점이다. 오직 데이터가 가지는 자기 유사성(self-similarity)에 의해서 서로 다른 클러스터로 나뉘게 된다. 나누어진 클러스터에 대해 의미를 부여하는 것은 전적으로 데이터 마이너(data miner)에 달려 있다. 이는 고객 세그멘테이션 등에 사용된다.

이러한 군집화 기술은 전체 데이터의 분포상태나 패턴 등을 찾아내는데 유용하게 사용된다.



[그림 2-1] 클러스터링을 통한 군집 생성

[Fig. 2-1] Cluster Generation through Clustering

나) 연관 규칙(Association Rule)

연관 규칙은 장바구니 분석이라고도 하는 것으로서 구매행위에 있어서 특정 아이템과 다른 아이템 간에 어떤 연관 관계가 있는 지를 찾아보는 것이다. 즉, 항목 집합(item set)들로 이루어진 데이터베이스에서 항목이 동시에 나타나는 성향(항목 A와 B가 같이 나타나는 경향)에 대한 관계성을 표현한 것이다. 비록 다른 데이터 마이닝 기법에 비해 단순하지만 적용이 쉽고 상당히 의미 있는 정보를 제공해 준다는 측면에서 인터넷 쇼핑몰의 상품 추천 시스템에서 널리 적용되는 기법중의 하나이다[6].

적용사례로는 미국의 대형 편의점의 소비자 구매 데이터에 이 기술을 적용한 결과 '아기 일회

용 기저귀를 사는 사람은 맥주도 같이 구매 한다'는 연관규칙을 발견한 것이다. 이를 마케팅에 활용하기 위해 맥주 앞에 기저귀를 진열하고, 기저귀를 할인하여 판매함으로써 맥주의 판매량을 늘리는 전략을 실행할 수 있었다.

그러나 연관 규칙을 이용한 추천 기법은 개인화된 추천이 아니며, 상품의 카테고리구조나, 시간적인 고객의 관심도를 반영한 추천도 아니다.

다) 순차 패턴(Sequential Pattern)

대용량의 데이터베이스로부터 순차패턴(Sequential Pattern)을 발견하는 문제는 지식 발견(Knowledge discovery)또는 데이터 마이닝(Data mining)분야에서 중요한 패턴 추출 문제로 알려져 있다.

순차패턴이란 트랜잭션 안에서 발생된 항목들 간의 연관규칙에 시간의 변이를 추가한 지식의 한 형태이며 데이터 마이닝의 여러 기법들 중 하나로서 연관규칙과 유사하지만 시퀀스로 구성된 주어진 데이터베이스에서 빈번히 나타나는 시퀀스를 찾는 방법이다[6][7]. 다시 말해 시간적인 순서에 따른 페이지의 집합을 나열했을 때 세션 간의 아이템의 연관성을 찾아 주는 데이터 마이닝 기술으로써 이는 한 요인에 영향을 미치는 선행 요인을 찾아내는 것이다[8].

이러한 순차패턴은 1995년 Agrawal에 의해 처음 소개된 이후 고객의 구매패턴 분석을 통한 구매 예측은 물론 의료 분야에서는 질병 발생 순서패턴에 따른 진료 및 투약에 관한 귀중한 정보를 제공하여 왔다[6][9]. 뿐만 아니라 최근에는 웹 로그를 통해 사용자가 웹 사이트에 접속해서 방문하는 경로(user web navigation path)를 통해서 사이트 구성과 광고배치를 결정하는 등 적용 범위가 다양화 되어 가고 있다[3][10].

3. 상품 추천 시스템 설계

3.1 상품 추천 시스템 정의

본 논문에서는 대용량 계층구조와 스트림데이터를 처리하기 위한 대형 쇼핑몰에서의 상품 추천 시스템을 모델링 한다[8][11].

가) 데이터 모델링

사용자 집합은 개개의 사용자를 u 라 정의하고 쇼핑몰을 이용한 모든 사용자들은 집합 U 로 표현한다.

$$U = (u_1, u_2, u_3, \dots, u_n) \quad (1)$$

상품 집합은 개개의 상품을 p 라 정의하고 쇼핑몰의 모든 상품들은 집합 P 로 표현한다.

$$P = (p_1, p_2, p_3, \dots, p_m) \quad (2)$$

상품의 카테고리는 c 라 정의하고 전체 카테고리 집합 C 는 다음과 같이 구성되어 있다.

$$C = (c_1, c_2, c_3, \dots, c_o) \quad (3)$$

상품 카테고리는 트리 구조를 갖는다.

$$T: C \rightarrow C \quad (4)$$

여기서, $T(C_c) = C_p$ 이 경우, C_p 는 C_c 의 부모 카테고리 이고 C_c 는 C_p 의 자식 카테고리이다.

나) 웹 로그 데이터 모델링

사용자의 상품과 카테고리에 대한 접근은 다음과 같은 형태로 웹 로그나 데이터베이스에 기록된다.

user	category	product	access-time
-------------	-----------------	----------------	--------------------

[그림 3-2] 웹 로그 포맷

[Fig. 3-2] Web Log Format

이러한 형태의 웹 로그 데이터는 사용자의 상품에 대한 관심도 정보를 추가하여 배치 작업을 통해 일 단위로 아래와 같은 형태로 가공할 수 있다. 여기서 count는 사용자의 상품 혹은 상품 카테고리에 대한 관심도(preference)를 의미한다.

user	category	date	count
user	product	date	count

[그림 3-3] 가공된 웹 로그 데이터

[Fig. 3-3] Modified web Log Data

이때, 관심도 함수 $\text{Pref_CD}(u, c_k, d)$ 는 사용자 u 의 날짜 d 에 상품 카테고리 c_k 에 대한 관심도로 정의한다.

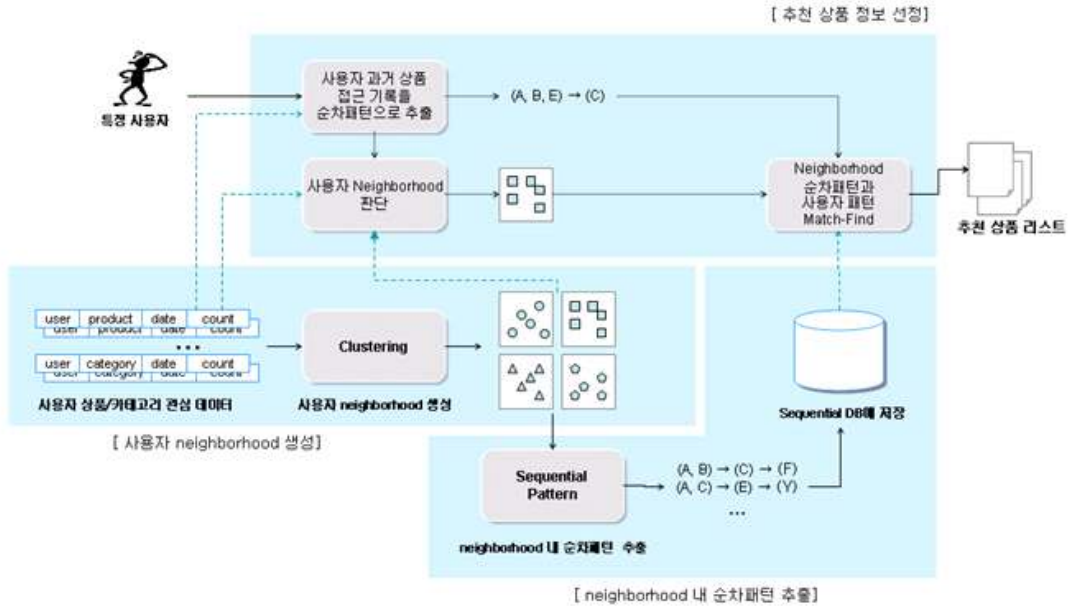
다) 상품 추천 방법 모델링

상품 추천 시스템에서 추천 상품 리스트를 생성하기 위한 방법을 다음과 같다.

첫째, 사용자 neighborhood를 생성한다. 사용자 neighborhood란 특정 사용자와 유사한 상품 선호도를 갖는 사용자 집합으로 웹 로그 정보 중 사용자들의 관심 상품 및 관심 카테고리에 정보를 기반으로 클러스터링 기법을 이용하여 계산한다.

둘째, 생성한 사용자 neighborhood 내에서 상품 접근 순차 패턴을 추출한다. 이 방법은 각 neighborhood에 소속된 사용자들의 상품 접근 웹 로그 데이터에서 순차 패턴을 추출한 것으로 일 단위로 로그 데이터를 가공하여 데이터베이스에 준비한다.

셋째, 추천 시스템에 사용자가 주어지면 사용자의 과거 상품 접근 기록을 이용해 상품 접근 순차를 생성하고, 사용자의 neighborhood로부터 추출한 순차패턴 정보를 검색하여 사용자의 순차를 포함하는 순차패턴이 있는지 Match-Find알고리즘을 이용해 찾게 된다. 순차 패턴이 존재하면 패턴태의 상품을 추천 상품으로 선정하게 된다.



[그림 3-5] 추천 방법 모델링

[Fig. 3-5] Modeling of Recommendation method

3.2 상품 추천 시스템 알고리즘

가) 정의와 규칙

추천은 주어진 사용자에게 대해 추천 상품을 선정해 주는 것이다. 추천은 위에서 설명한 사용자 neighborhood와 neighborhood별 순차 패턴 추출 정보를 이용해 이루어지게 된다.

주어진 사용자에게 대해 추천 상품을 선정하는 과정은 다음과 같다.

첫째, 사용자 u 의 과거 상품 접근 기록을 순차 패턴으로 표현한다.

$$S_u = S_{u1} \rightarrow S_{u2} \rightarrow \dots \rightarrow S_{un} \quad (5)$$

둘째, 사용자 u 가 속한 클러스터 I 를 판단한다.

셋째, 클러스터 I 에서 사용자 u 의 순차패턴 S_u 를 닮은 순차패턴을 찾아 추천 상품을 결정한다.

이 처리 과정에서 순차 패턴 S_u 를 닮은 순차패턴을 찾는 과정이 Match-Find 알고리즘 이다.

나) Match-Find 알고리즘

Match-Find 알고리즘에서는 주어진 sequential DB 내에서 사용자의 순차패턴을 포함하는 순차패턴을 찾는다. 즉, $\text{include}(\text{Head}(s), S_u)$ 를 만족하는 순차패턴의 $\text{tail}(s)$ 를 추천해주면 된다. 이러한 $\text{tail}(s)$ 가 복수 개 일 수도 있으므로 지지도(support)를 기준으로 ordering 하여 반환해준다. 단순화된 모델을 위해 클러스터 내 sequential DB에서 데이터를 추출할 때 tail 사이즈가 1인 순차패턴만을 추출한다.

알고리즘1. Match-Find 알고리즘
입력: user u 출력: Product P begin 1. Get the user's product access sequence S_u 2. Get the user's cluster number i 3. For each sequence S_i in SeqDB_i if $\text{include}(\text{Head}(s), S_u)$ insert $\langle \text{Tail}(S_i), \text{support } S_i \rangle$ to Ret 4. Order Ret by support, return top p product end.

[그림 3-6] Match-Find 알고리즘

[Fig. 3-6] Match-Find Algorithm

4. 상품 추천 시스템 구현

4.1 시스템 Framework 구조

본 논문에서는 실시간 추천을 위한 구조로 반영하기 위하여 학습 모듈과 추천 모듈로 분리하였다.



[그림 4-1] 시스템 구조

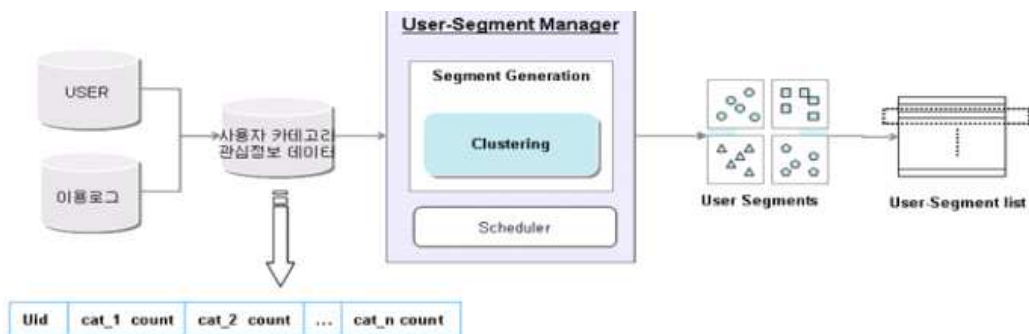
[Fig. 4-1] System Structure

4.2 학습 모듈과 추천 모듈

학습 모듈은 일별 배치작업으로 수행된다. 웹 로그로부터 사용자의 상품/카테고리 관심 데이터를 생성하고, 생성된 사용자의 상품/카테고리 관심 데이터를 기반으로 neighborhood를 미리 계산한다. 클러스터링 기법으로 얻은 neighborhood 단위로 순차 패턴마이닝을 통해 클러스터별 순차패턴을 추출한다.

4.3 학습 모듈과 추천 모듈

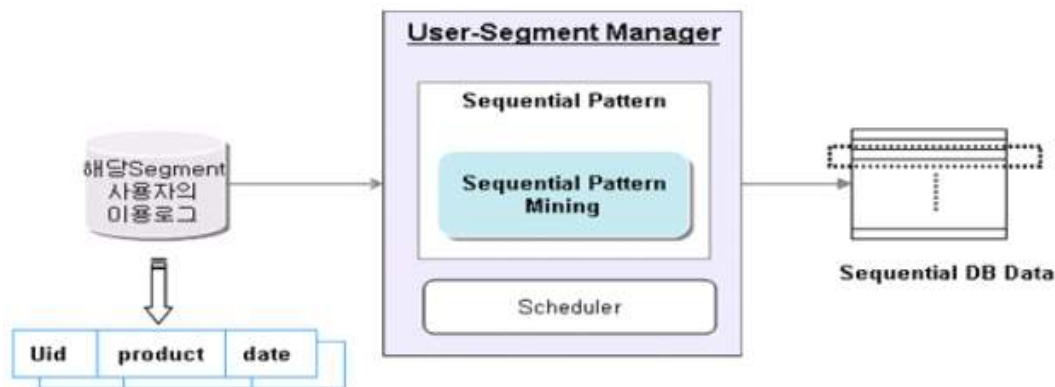
학습 모듈은 일별 배치작업으로 수행된다. 웹 로그로부터 사용자의 상품/카테고리 관심 데이터를 생성하고, 생성된 사용자의 상품/카테고리 관심 데이터를 기반으로 neighborhood를 미리 계산한다. 클러스터링 기법으로 얻은 neighborhood 단위로 순차 패턴마이닝을 통해 클러스터별 순차패턴을 추출한다.



[그림 4-3] 사용자 neighborhood 계산 과정

[Fig. 4-3] Counting Process of user's neighborhood

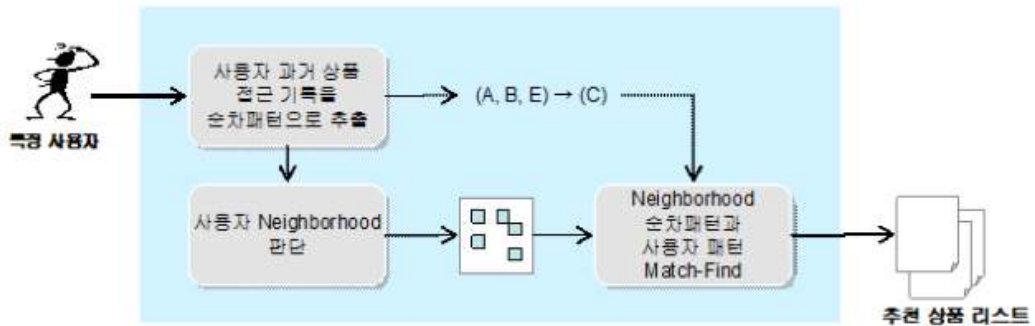
neighborhood 계산 방법은 K-means 클러스터링 기법을 사요하고 클러스터링이 끝나면 순차 패턴마이닝에 의해서 순차 패턴을 추출한다.



[그림 4-4] 클러스터 내 사용자들의 순차 패턴 추출

[Fig. 4-4] users sequence pattern extraction in cluster

추천 모듈은 실시간으로 처리되는 작업이다. 학습 모듈에서 생성된 정보를 이용해 실시간 추천을 수행한다. 특정 사용자의 과거 상품 접근 기록과 사용자가 속한 클러스터 내의 순차 패턴을 Match-Find 알고리즘을 통해 검색하고 그 결과로 추천 상품 리스트를 제공한다.



[그림 4-5] 추천 모듈에서의 작업 처리 과정

[Fig. 4-5] work flow process of recommendation module

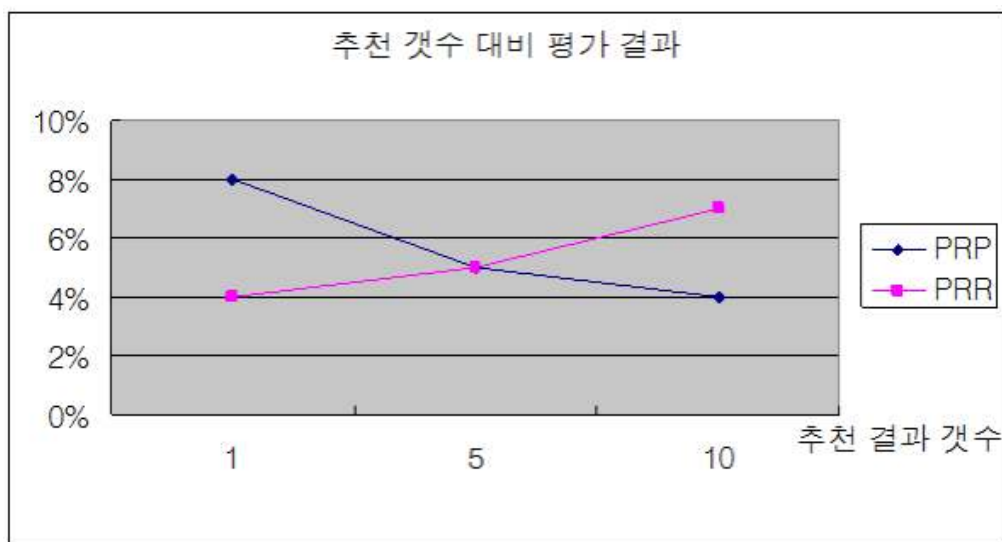
본 논문에서 구현한 K-means 알고리즘은 대용량 처리를 위해 마이닝 대상 데이터를 모두 디스크 상에 유지한 채 수행되도록 구현했고 초기 클러스터 수는 5로 주어진 경우 50만 사용자에게 대한 neighbor 계산이 10분 이하에 수행된다. 또한, 제안된 방법인 Match-Find 알고리즘은 전체 상품을 검색하는 방식이 아니라, 학습 단계에서 미리 발견된 패턴에서 검색하는 것으로서, 훨씬 빠른 연산이며, 학습된 패턴 수에 Linear한 알고리즘이다.

이와 같은 특성이 있으면서도, 개인의 과거 선호 상품패턴 정보를 활용하는 개인화된 추천 방식으로 추천 결과는 시계열적인 과거 패턴을 Match-Find 알고리즘에 의해서 반영하였다.

5. 실험 및 평가

상품 추천 시스템의 분석·평가 결과는 다음과 같다.

먼저 학습시간은 neighborhood 계산(= clustering 작업)은 10분이고 순차패턴 추출 시간은 10분/클러스터 이다. 그리고 추천 시간(= Match-Find)은 사용자당 1초가 걸렸다. 학습이 완료된 후, 사용자 10,000명에 대해 Predictive recommend precision(PRP)와 predictive recommend recall(PRR)을 구해 본 결과는 아래와 같다.



[그림 5-3] 시스템 평가 결과

[Fig. 5-3] The Result of system evaluation

추천 결과 개수를 증가 시키면, PRR은 증가하고 PRP는 감소하는 경향이 있다. 현실적으로는 추천 개수는 5개미만으로 제공되어야 하며, 신규 콘텐츠 등 다른 추천 방법들과 동시에 사용자에게 노출되기 때문에, 추천 개수의 제한이 많은 편이다.

PRP, PRR은 예측된 recall, 예측된 precision 이므로 낮은 숫자가 나온 것이고, 실제 사용자에게 노출된 후 feedback을 계산한 본래 의미의 recall, precision은 실험값 보다 높은 값이 나올 것으로 예상된다.

추가적인 실험 방법으로는 학습 기간과 예측 대상 기간을 변경하는 방법, 사용자 cluster 개수를 조절해보는 방법, 순차패턴 추출 시 지지도를 조절해보는 방법 등이 있을 수 있다.

6. 결론

본 논문에서는 상품수가 수십만 이상이고 고객수가 수백만 이상인 환경에 적용될 수 있고, 매일 변화하는 고객의 관심도를 반영하고 카테고리화 상품 정보를 활용하는 추천 시스템을 제안했다.

제안된 시스템은 매일 변화하는 고객의 관심도를 반영하기 위해, 매일 고객을 재 그룹화 하고 실시간 추천에 필요한 배치 작업을 수행한다. 고객 군 그룹과 작업은 K-means 군집화 알고리즘을 사용하고, 각 고객 그룹별로 관심 상품의 시계열적인 변화를 순차패턴 알고리즘을 사용하여 추출해둔다. 실시간 추천은 고객별 과거 관심 상품 내역과 배치 작업으로 추출한 순차 패턴을 비교하여 추천 상품을 결정한다. 이를 통해 대용량의 계층 구조를 가지면서 고객의 상품에 대한 관심 정보가 일 단위로 연속해서 변화하는 steam data 환경에 대한 새로운 추천 프레임워크를 제시했으며, 대형 쇼핑몰의 실제 데이터를 활용하여 그 효용성을 보였다.

향후 연구분야는 관심 있었던 상품 및 카테고리 내의 속성 등에 대한 학습을 고려하는 분야와 상품의 인기도나 신규정보 등을 고려하는 것이다.

References

- [1] Gartner group, website.(2000). "Personalization strategies / technologies", Conference presentation, SPGI Website 400 Cabrams.
- [2] S. J.. Ben, J. Konstan, J. Riedl, Data Mining and Knowledge Discovery. (2001), Vol.5, No.1, pp.115-153.
- [3] A. Asim, S. Essegaier, R. Kohli, Journal of Marketing Research. (2000), Vol.37, pp.363-375.
- [4] Y. H. Cho, J. K. Kim, and S. H. Kim, Expert Systems with Applications. (2002), Vol.23, No.3, pp.329-342.
- [5] Dragon, V. Richard,(1997). "Recommendation Systems-Advice From the web," PC Magazine, September 9, 1997.
- [6] Balabanovic, M, Shohm, Y. (1997). "Fab: Content-Based, Collaborative Recommendation", Communications of ACM, March. 1997.
- [7] B. W. Seo, Evolution of content recommendation algorithm, Broadcasting Trend & Insight, (2013).
- [8] K. N. Park and N. J. Kang, A Study on Development of Products Recommendation System based on e-CRM using Customer Equity, The Journal of Internet Electronic Commerce Research, Vol. 10, No. 1, (2010), pp.125-144.
- [9] Y. H. Cho and I. W. Kim, Predicting the Performance of Recommender Systems through Social Network Analysis and Artificial Neural Network, Journal of Intelligence and Information Systems, (2011), Vol. 16, No. 4.
- [10] S. H. Kwon, S. C. Lee and S. W. Kim, A Recommendation System Based on Message Passing, The conference of Korean Institute of Information Scientist and Engineers, (2011), Vol. 38, No. 2.
- [11] Y. J. Lee and K. J. Lee, Product Recommender Systems using Multi-Model Ensemble Technique, Journal of Intelligence and Information Systems, (2013), Vol. 19, No. 2.