

박사학위논문

온라인 광고 시장에서 예산 분배와
클릭 확률 예측

2019년

한 성 대 학 교 대 학 원

경 제 부 동 산 학 과

경 제 학 전 공

해 현 용

박사학위논문
지도교수 김상봉

온라인 광고 시장에서 예산 분배와 클릭 확률 예측

Estimation of Click Probability and Budget
Allocation in the Online Advertising Market

2018년 12월 일

한 성 대 학 교 대 학 원

경 제 부 동 산 학 과

경 제 학 전 공

해 현 용

박사학위논문
지도교수 김상봉

온라인 광고 시장에서 예산 분배와 클릭 확률 예측

Estimation of Click Probability and Budget
Allocation in the Online Advertising Market

위 논문을 경제학 박사학위 논문으로 제출함

2018년 12월 일

한 성 대 학 교 대 학 원

경 제 부 동 산 학 과

경 제 학 전 공

해 현 용

해현용의 경제학 박사학위 논문을 인준함

2018년 12월 일

심사위원장 _____(인)

심 사 위 원 _____(인)

심 사 위 원 _____(인)

심 사 위 원 _____(인)

심 사 위 원 _____(인)

국 문 초 록

온라인 광고 시장에서 예산 분배와 클릭 확률 예측

한 성 대 학 교 대 학 원
경 제 부 동 산 학 과
경 제 학 전 공
해 현 용

광고란 서비스를 구독하거나 상품을 구매하려는 잠재적인 고객을 대상으로 하는 마케팅 기법이다. 광고는 또한 브랜드와 관련된 광고의 반복된 노출을 통해 브랜드 이미지를 만드는 방법이기도 하다. 기술의 발전과 함께 사람들의 생활 반경은 오프라인에서 온라인으로 옮겨졌다. 온라인 광고 산업의 핵심 목표는 사용자 경험을 바탕으로 사용자의 인터넷 경험을 향상시키고, 광고주의 광고 효과를 높이며, 매체의 이윤을 극대화하는 것이다. 이를 위해 쿠키를 이용하여, 브라우저의 행동 정보를 수집한다. 그리고 수집된 정보를 바탕으로 브라우저의 사용자 프로파일링을 한다. 그런 다음 프로파일링 한 정보를 바탕으로 개인화된 광고를 내보낸다. 이런 개인화 광고를 내보내기 위해서 시스템이 필요해졌다.

본 논문은 이 온라인 광고 시스템에서 알고리즘에 관한 연구이다. 먼저 온라인 광고 시장에 대한 개괄적인 소개를 한다. 그 후 온라인 광고 시장의 경제학적 특성을 논의하고, 주요 문제를 제시한다. 온라인 광고 시장은 그 시장 규모나 발달 속도에 비해 우리나라에 알려진 바가 적고, 해외에서는

활발히 연구되고 있다. 본 연구는 온라인 광고 시장에 핵심 문제들에 대한 최신 알고리즘에 대한 개괄적인 소개와 더불어 개선할 수 있는 방향을 제시하는 방안을 제안함으로써 계량경제학 분야에 학술적인 기여를 할 수 있을 것으로 기대한다.

첫 번째로 온라인 광고 시장의 주요 문제인 클릭 확률 추정과 관련된 알고리즘 소개와 개선 방안에 대해 논의한다. 본 연구에서는 클릭 확률 추정 모델을 만들 때에 베이스 모델로는 FM-FTRL을 사용했다. 또한 추가적인 변수로 입력 벡터의 비선형 그룹을 찾아내는 K-평균 클러스터링 알고리즘의 예측치를 사용했다. 그리고 그라디언트 부스팅 모델의 잎 노드들을 추가적인 변수로 사용했다. K-MEANS 변수와 GBDT 변수가 각각 원 자료의 비선형 특성을 전체 자료의 10%만 가지고도 잘 포착하며, 일반화된다고 볼 수 있다.

그 다음으로 온라인 광고 시장에서 또 하나의 중요한 문제인 예산 분배에 대해 논의한다. 광고주는 시간에 따라서 부드럽게 게재되는 예산 할당을 추구한다. 시간에 따라서 부드럽게 예산 분배가 된다면, 같은 광고비를 지속적으로 사용할 수 있고, 더 넓은 잠재고객에게 도달할 수 있기 때문이다. DSP는 개별 증권사와 매우 유사하다. DSP를 마치 증권사처럼 생각한다면, 예산 분배는 포트폴리오를 작성하는 것이라 생각할 수 있다. 광범위한 주식 중에서 최적 기대 이윤을 가져다 줄 최적 포트폴리오를 결정하는 문제와 동일시 될 수 있다. 일반적인 예산 분배 방법은 스로틀링과 입찰 수정 방법이 있다. 본 연구에서는 스로틀링 방식을 사용했다. 개선점으로는 DSP 입장에서 포트폴리오 이론을 바탕으로 한 최적 입찰 확률을 추정하였다. Sharpe Ratio를 활용해 리스크 대비 투자 결과를 측정하고 최대 비율이 되는 가중치를 사용하여, 예산 소진을 거의 다 하면서도 퍼포먼스를 향상시키는 방안을 제시하였다.

온라인 광고 시장은 빅 데이터라 불리는 방대한 양의 데이터를 가지기 때문에 스몰 데이터에서는 문제가 되지 않던 것들이 문제가 되는 경우가 많다. 단순히 양이 많아짐에 따라 발생하는 문제들을 처리하는 방법 또한 잘 알려져 있지 않고, 구전되는 경우가 많기 때문에 해당 내용에 대한 논의를 통해 실무에 도움이 되기를 기대한다.

【주요어】 대규모 자료, 실시간 학습, 온라인 광고, 스택킹 기법, 예산 분배,
클릭 확률 추정

목 차

I. 서 론	1
1.1 연구의 목적	1
1.2 연구의 범위와 방법 및 기여	5
1.2.1 연구의 범위	5
1.2.2 연구의 방법	5
1.2.3 연구의 기여	6
II. 온라인 광고 시장	7
2.1 온라인 광고 시장에 대한 배경	7
2.2 온라인 광고 시장의 경제학적 특성	15
2.2.1 두 번째 가격 경매	16
III. 벤치마킹 자료	19
3.1 자료 소개	19
3.2 기초 통계량	25
3.3 자료 특성	28
IV. 클릭 확률 추정	36
4.1 서론	36
4.2 이론적 배경	38
4.2.1 로지스틱 회귀 (Logistic Regression)	38
4.2.2 선행 정규화를 따르는 로지스틱 회귀	44
4.2.3 FM (Factorization Machine)	45
4.2.4 의사결정 나무 (Decision Tree)	46
4.2.5 그라디언트 부스트 나무 모형 (Gradient Boosted Tree Model)	47

4.2.6 스택킹 앙상블 (Stacking Ensemble)	50
4.2.7 K-평균 클러스터링 (K-Means Clustering)	51
4.2.8 부호화 기법 (Hashing Trick)	51
4.2.9 결측 자료 보정	52
4.3 연구 설계	57
4.3.1 데이터와 문제 정의	57
4.3.2 분석 방법	58
4.4 결론	60
4.4.1 학습 결과	60
V. 예산 분배	69
5.1 서론	69
5.2 이론적 배경	72
5.2.1 부드러운 예산 분배	72
5.2.2 스마트 예산 분배	77
5.2.3 현대 포트폴리오 이론	79
5.3 연구 설계	80
5.4 결론	81
5.4.1 학습 결과	81
VI. 결 론	85
참 고 문 헌	88
ABSTRACT	92

표 목 차

[표 3-1] 광고주 산업	19
[표 3-2] iPinYou 제공 로그 형태 예시	23
[표 3-3] 학습 자료 기초 통계량 1	25
[표 3-4] 학습 자료 기초 통계량 2	26
[표 3-5] 평가 자료 기초 통계량 1	26
[표 3-6] 평가 자료 기초 통계량 2	27
[표 3-7] 광고주별 입찰 가격의 기초통계량	35
[표 3-8] 광고주별 낙찰 가격의 기초통계량	35
[표 4-1] 사용자 특성 벡터의 형태 예시	57
[표 4-2] FM-FTRL의 학습 결과	60
[표 4-3] FM-FTRL+K-MEANS의 학습 결과	61
[표 4-4] FM-FTRL+GBDT의 학습 결과	61
[표 4-5] FM-FTRL+K-MEANS+GBDT의 학습 결과	62
[표 4-6] 학습 결과 요약표	63
[표 5-1] 예산 분배의 결과 (클릭 수)	81
[표 5-2] 예산 분배의 결과 (예산소진율)	82
[표 5-3] 예산 분배의 결과 (기본 입찰가)	82
[표 5-4] 예산 분배의 결과 (CTR)	83
[표 5-5] 예산 분배의 요약표	83

그 림 목 차

[그림 1-1] 쿠키 동기화 방법	2
[그림 1-2] RTB(Real-Time Bidding) 시장 구조	3
[그림 1-3] 사용자 추적 방법	4
[그림 2-1] 초기 광고 형태	7
[그림 2-2] Ad Server 도입	8
[그림 2-3] Ad Network 등장	9
[그림 2-4] Ad Network 역할	10
[그림 2-5] Ad Network 사이의 인벤토리 교환	11
[그림 2-6] Ad Network 사이의 인벤토리 거래	12
[그림 2-7] Ad Exchange의 등장	13
[그림 2-8] 두 번째 가격 경매 예시	16
[그림 3-1] DSP, ADX, 유저 상호작용	20
[그림 3-2] 평가 프로토콜	21
[그림 3-3] 요일별 광고주별 CTR 차이	28
[그림 3-4] 시간대별 광고주별 CTR 차이	29
[그림 3-5] 시간대별 광고주별 CTR 차이 (2997 제외)	30
[그림 3-6] 지역별 광고주별 CTR 차이	30
[그림 3-7] ADX별 광고주별 CTR 차이	31
[그림 3-8] 소재 크기별 광고주별 CTR 차이	32
[그림 3-9] 학습 자료에서의 입찰가격 커널 밀도 함수	32
[그림 3-10] 학습 자료에서의 낙찰가격 커널 밀도 함수	33
[그림 3-11] 입찰 가격과 낙찰 가격의 결합 분포	34
[그림 4-1] 과-적합 (Over-fitting)	42
[그림 4-2] Lasso 와 Ridge에서 계수 값 추정	43
[그림 4-3] 의사결정나무의 예	47
[그림 4-4] 학습과 예측 사이의 자료 편향	53
[그림 4-5] 낙찰 확률 계산 예시	55

[그림 4-6] 1458 광고주의 ROC Curve	64
[그림 4-7] 2259 광고주의 ROC Curve	64
[그림 4-8] 2261 광고주의 ROC Curve	65
[그림 4-9] 2821 광고주의 ROC Curve	65
[그림 4-10] 2997 광고주의 ROC Curve	66
[그림 4-11] 3358 광고주의 ROC Curve	66
[그림 4-12] 3386 광고주의 ROC Curve	67
[그림 4-13] 3427 광고주의 ROC Curve	67
[그림 4-14] 3476 광고주의 ROC Curve	68
[그림 5-1] 확률적 조정과 입찰 수정 전략	69
[그림 5-2] 다양한 예산 분배 형태	70
[그림 5-3] 예측 CTR 별 광고 요청 수의 분포	75

I. 서 론

1.1 연구의 목적

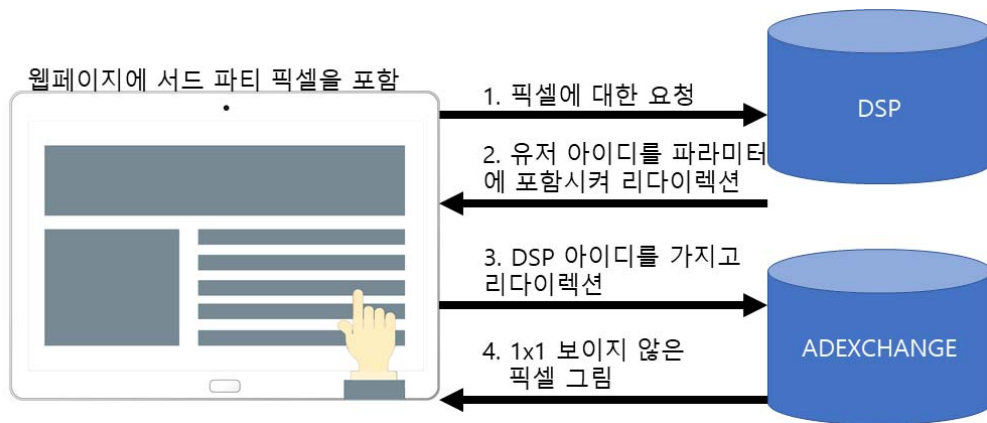
광고는 서비스를 구독하거나 상품을 구매하려는 잠재적인 고객을 대상으로 하는 마케팅 기법이다. 광고는 또한 미디어에서 브랜드와 관련된 광고의 반복된 노출을 통해 브랜드 이미지를 만드는 방법이다. 이전부터 있어 왔던 텔레비전, 라디오, 신문, 잡지 등은 전통적인 광고의 주요 채널이었다. 그러나 인터넷의 발달로 사람들은 정보의 대부분을 온라인에서 찾을 수 있게 되었다. 사람들은 온라인상에서 특정 웹사이트를 돌아다니고 전자 상거래 등의 거래 행위를 포함한 일련의 사회적 행동을 반복한다. 온라인상에서의 활동이 활발해짐에 따라 오프라인에서의 서비스 등이 온라인으로 이동하고 있고, 잠재적 고객에게 도달하려는 광고주는 자연스럽게 온라인에도 광고를 게재하고 싶어 하는 욕구가 생기게 되었다.

온라인 광고 산업은 IT 산업에서 가장 빠르게 진보하는 영역이다. 웹 디스플레이 광고에서부터 모바일 광고까지 최신 기술에는 언제나 광고가 접목되고 있다. 온라인 광고 산업의 핵심 목표는 사용자 경험을 바탕으로 한 개인화 추천을 통해 사용자의 인터넷 경험을 향상시키고, 광고주의 광고 효과를 높이며, 매체의 이윤을 극대화하는 것이다. 결국 온라인 광고 산업은 사용자(Audience, User), 광고주(Advertiser), 매체(Publisher)가 활동하는 시장이 된다.

그 중 디스플레이 광고란 웹 배너의 형태로 웹 페이지에 나타나는 광고이다. 주로 브랜딩 용도로 사용되고, 쿠키와 브라우저 히스토리를 이용하여 사용자의 데모그래픽이나 관심사를 파악한다. 배너 광고는 파악한 프로파일링 정보를 이용해 오디언스에게 다양한 방식으로 타게팅 된다. 행동기반 리타게팅, 데모그래픽 타게팅, 위치 기반 타게팅, 관심사 타게팅, 사이트 기반 타게팅, 유사 유저 타게팅 등이 주로 사용되는 방법이다. 일반적으로 배너는 이미지 또는 간단한 사운드와 영상이 포함된 형태로

제공된다. 특정 페이지의 트래픽을 광고주의 홈페이지로 전이시키는 것이 기본 기능이며, 특정 페이지의 유저들에게 특정 상품과 브랜드를 인지시키기 위한 목적으로 주로 사용된다. 일반적으로 지불 방식에는 클릭 당 비용(CPC, Cost Per Click) 방식 등이 있다.

쿠키(Cookie)란 HTTP 웹 문서의 일부이며, 사용자가 브라우저를 통해 특정 웹사이트를 방문할 경우 그 사이트의 도메인 이름에 종속되어 작성되는 작은 크기의 기록 정보 파일을 말한다. 쿠키는 사이트의 로그인 정보, 방문 상세 페이지 혹은 페이지 상품 정보 등의 다양한 정보를 포함할 수 있으며, 각 도메인 이름에 종속되기 때문에 다른 사이트의 쿠키는 해당 사이트의 정보를 수집할 수 없다. 그러나 추후에 서술할 온라인 광고 회사에서는 매체 사이트 등에 자신들의 사이트를 경유하는 아주 작은 1x1 픽셀의 이미지 주소를 심어두고, 호출 요청을 통해 실질적으로 광고 사이트에도 접속한 것과 같은 형태로 사용자의 쿠키 정보를 수집한다.

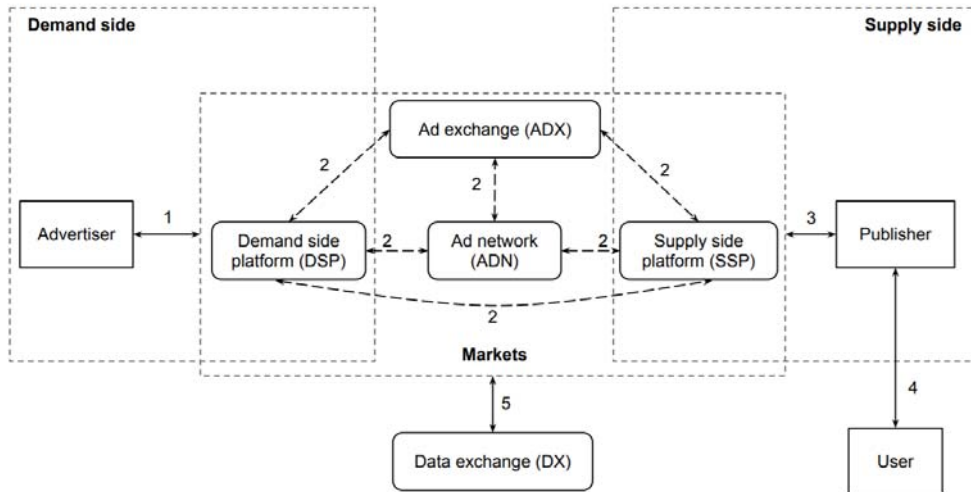


[그림 1-1] 쿠키 동기화 방법

결국 온라인 광고 시장은 각 개별 사용자에게 대한 비-식별 처리된 데이터인 쿠키데이터를 통해서 각 개인에게 직접적이고 통합적으로 접근할 수 있다. 후술할 온라인 광고 시장은 완전경쟁시장에 비견되기 때문에 완전경쟁시장에서의 개별 주체의 행동에 대한 분석이 가능하다.

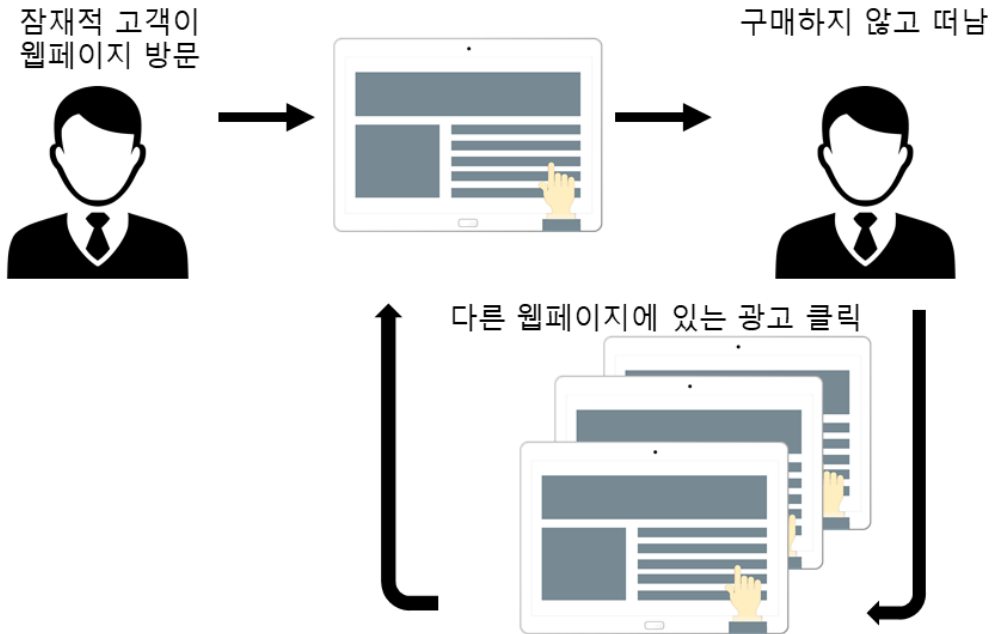
본 논문에서는 광고 산업과 광고 시장의 경제학적 특성을 소개한다. 그런

다음 광고 산업에서의 핵심적인 문제를 정의하고, 해당 문제를 해결하기 위해 경제학적 접근을 바탕으로 해결 방안을 제시하는 것을 목적으로 한다.



[그림 1-2] RTB(Real-Time Bidding) 시장 구조

자료: Wang, J., Zhang, W., & Yuan, S. (2016)



[그림 1-3] 사용자 추적 방법

사용자를 추적하는 방법은 먼저 위에 쿠키 동기화 방법에서 각 광고 수요자와 공급자의 서로 1x1 픽셀의 그림을 참조함으로써 아이디를 동기화 시킨다. 그런 후에는 이제 잠재적 고객의 행동을 네트워크 내에서 추적할 수 있고, 잠재적 고객이 다시 매체에 방문 했을 경우, 해당 광고 수요자의 광고를 노출시킬 수 있다.

1.2 연구의 범위와 방법 및 기여

1.2.1 연구의 범위

본 연구의 내용적 범위는 온라인 광고 시장에서의 클릭 확률 추정과 예산 분배이다. 클릭 확률 추정에서는 클릭 확률은 광고 요청 한 건에 대한 클릭 확률로 한정한다. 예산 분배에서는 평가 기간 동안의 각 개별 캠페인에 대한 예산 분배로 한정한다.

시간적 범위는 클릭 확률의 경우 논의가 온라인 학습이 진행되는 것이기 때문에 실시간이다. 예산 분배의 경우는 평가 기간 동안의 각 개별 캠페인에 대한 예산 분배를 논의함으로 각 캠페인별로 다를 수 있다.

본 연구의 데이터 상 범위는 iPinYou DSP(Demand Side Platform)에서 벤치마킹용으로 제공한 데이터이다. 본 연구에 활용된 데이터의 정합성은 제공 기관의 신뢰성으로 같음한다. 본 연구에서 논의될 문제들에 대한 알고리즘은 범용으로 사용할 수 있으므로 각 구현될 데이터에 따라 서로 다른 효과를 가질 수 있다.

1.2.2 연구의 방법

본 연구는 총 6개의 장으로 이루어져 있다. 1장은 서론으로 논문의 배경이 되는 산업에 대한 요약과 연구를 진행하게 된 계기를 서술하였다. 2장에서는 온라인 광고 산업에 대한 배경과 온라인 광고 시장의 경제학적 특성 등을 자세히 설명한다. 3장에서는 클릭확률 추정과 예산 분배를 학습하고 평가하기 위한 벤치마킹 자료에 대해 소개한다. 4장에서는 클릭 확률 추정의 필요성을 논의하고, 추정 방법론과 실제 구현 시에 고려해야 할 사항 등을 서술한다. 5장에서는 온라인 광고 경매에서 입찰 시에 예산 분배에 대한 내용을 서술한다. 6장에서는 종합적으로 요약해 결론을 내는 것으로 마친다.

1.2.3 연구의 기여

본 논문은 온라인 광고 시장에서 경제학을 활용하는 것을 주제로 하여, 핵심문제들에 대한 종합적인 분석을 진행한 연구이다. 온라인 광고 시장은 그 시장 규모나 발달 속도에 비해 우리나라에 알려진 바가 적고, 해외에서는 활발히 연구되고 있다. 본 연구는 온라인 광고 시장에 핵심 문제들에 대한 최신 알고리즘에 대한 개괄적인 소개와 더불어 개선할 수 있는 방향을 제시하는 방안을 제안함으로써 계량경제학 분야에 학술적인 기여를 할 수 있을 것으로 기대한다. 본 연구는 두 가지 온라인 광고 시장의 핵심적인 문제에 대해 다룬다. 첫째, 클릭률 추정에서는 FM-FTRL(Factorization Machine-Follow The Regularization Leader) 알고리즘을 기반으로 GBM(Gradient Boosting Model)의 잎 노드와 K-평균 변수를 사용하는 방법을 사용하여 성능을 높인다. 둘째, 예산 분배 알고리즘에 있어서 기존 방식과는 달리 현대 포트폴리오 이론을 적용하여 DSP(Demand Side Platform)의 기대 효과를 높일 수 있는 포트폴리오를 제안한다. 온라인 광고 시장은 빅 데이터라 불리는 방대한 양의 데이터를 가지기 때문에 스몰 데이터에서는 문제가 되지 않던 것들이 문제가 되는 경우가 많다. 단순히 양이 많아짐에 따라 발생하는 문제들을 처리하는 방법 또한 잘 알려져 있지 않고, 구전되는 경우가 많기 때문에 해당 내용에 대한 논의를 통해 실무에 도움이 되기를 기대한다.

Ⅱ. 온라인 광고 시장

2.1 온라인 광고 시장에 대한 배경

본 절은 이현채(2012)의 글을 바탕으로 배경을 요약 정리하였다. 온라인 광고 시장은 광고주(Advertiser)와 매체(Publisher), 그리고 사용자(User, Audience)로 구성된다. 인터넷 초기에는 홈페이지를 보유하고 있는 매체와 광고주만이 있던 시기였다. 광고주는 인터넷의 발달에 따라 오프라인에서 온라인으로 그 활동 반경을 넓히고 있는 잠재적 고객인 사용자에게 광고를 노출하고 싶은 욕구가 있었다. 광고를 하고자 하는 광고주는 원하는 매체에 연락을 하여 계약 조건을 합의하고, 배너 광고의 형태와 소재를 전달한다. 매체는 해당 자료를 받아 HTML 페이지 수정을 통해, 배너 이미지를 부착하여 홈페이지를 만든다.

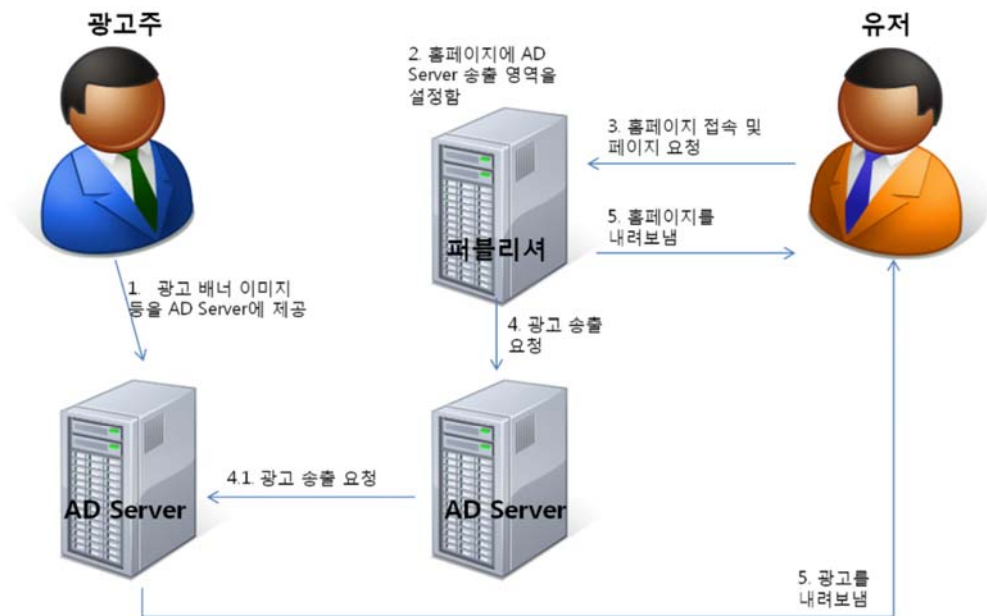


[그림 2-1] 초기 광고 형태

자료: 이현채. (2012년 1월 11일), [Ads. Study] Display Ads. 3편 - 초기 광고의 거래(?) 형태와 진화, 검색일 : 2018년 10월 20일,
<http://itmobile.tistory.com/entry/AdsStudyDisplayAds>

이러한 방식은 매체가 항상 배너를 받아 수정해야하는 비효율성이 있다. 또한 광고주 입장에서든 개별 매체에 배너가 제대로 붙어있는지 확인해야 했기 때문에 관리의 어려움이 있다. 따라서 그 이후에 Ad Server가

도입되었다. Ad Server는 광고주로부터 배너 이미지 등을 받아 저장하고 있고, 매체는 HTML 내에 특정 영역(Area, Inventory)을 지정을 해두고, 해당 영역은 Ad Server와 연결 되어 있다. 이제 사용자가 매체 홈페이지에 접속 하면, 사용자에게 홈페이지를 보여주고, 동시에 매체는 Ad Server에 광고 송출 요청을 보내면 Ad Server는 해당 영역에 내보내야 하는 배너를 사용자에게 보여주게 된다. Ad Server는 광고주 쪽과 매체 쪽이 있으며, 광고주 쪽 광고주에 대한 배너를 관리하고 통계 내며 수수료를 얻고, 매체 쪽은 광고 영역을 관리하며 수수료를 얻는다.

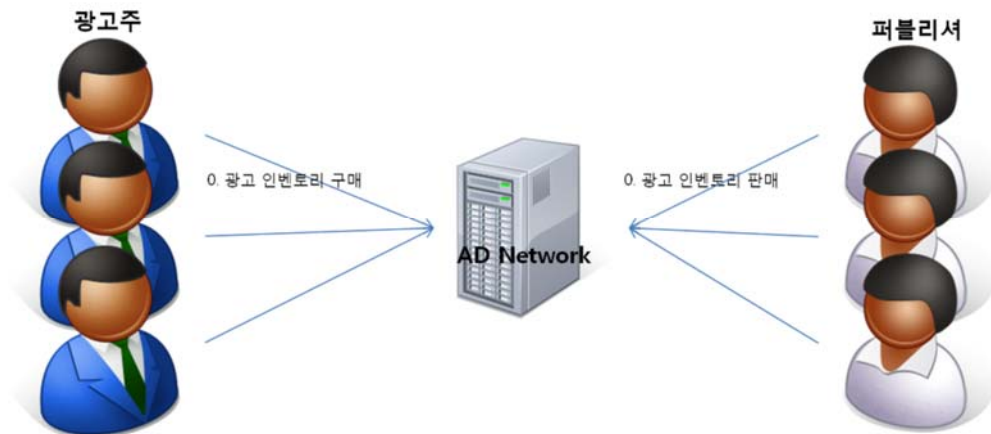


[그림 2-2] Ad Server 도입

자료: 이현채. (2012년 1월 11일), [Ads. Study] Display Ads. 3편 - 초기 광고의 거래(?) 형태와 진화, 검색일 : 2018년 10월 20일,
<http://itmobile.tistory.com/entry/AdsStudyDisplayAds>

매체는 광고 영역을 판매 혹은 제공 할 수 있고, 광고주는 광고 영역을 구매한다. 온라인 광고 시장에서의 수요자는 광고주이며, 공급자는 매체가 된다. 판매되는 상품의 종류는 특정 영역에 광고를 송출 할 수 있는 횟수가 된다. 이런 상품은 비행기 티켓이나 극장 티켓처럼 시간 소멸성 상품이다.

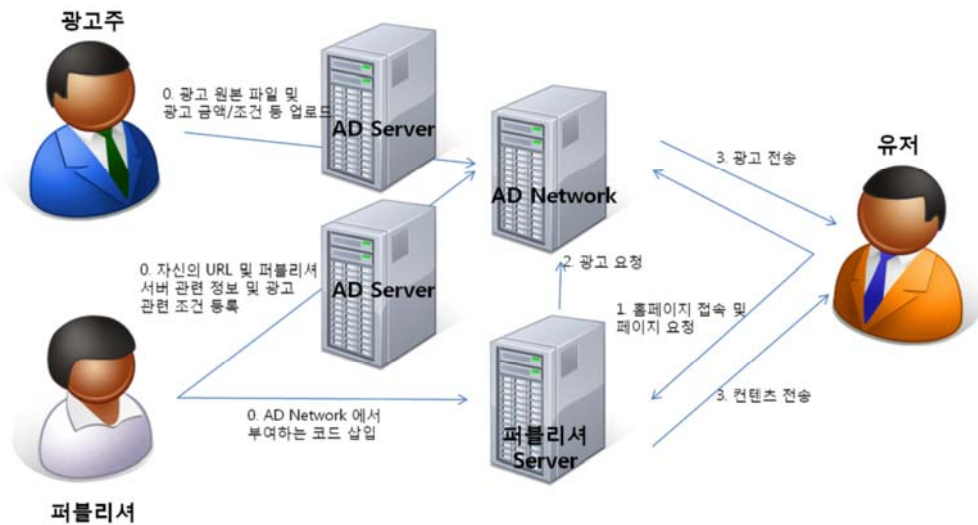
시간에 팔리지 않으면 판매 기회를 잃어버리게 된다. 말하자면 사용자가 매체 홈페이지에 접속했다 하더라도, 광고를 노출할 것이 없으면 해당 기회를 잃어버리게 되는 것이다. 그렇기 때문에 매체는 독점적 경쟁을 하는 공급자처럼 행동하며, 스스로의 수익을 극대화하기 위해 상품의 공급량과 가격 등을 조정한다. 광고주는 잠재적 고객인 사용자에게 광고를 송출함으로써, 자신의 브랜드 이미지를 확립하고자하는 욕구에 따라 광고 매체에 돈을 지불하고 광고 송출 기회를 구매한다. Ad Server의 등장으로 광고주와 매체는 효율적인 거래를 할 수 있었지만, 광고주의 수가 늘고, 매체의 수도 점차 늘면서, 각 Ad Server를 통합하는 거래 대행사가 생겨나게 되었다.



[그림 2-3] Ad Network 등장

자료: 이현채. (2012년 1월 18일), [Ads. Study] Display Ads. 4편 - AD Network의 역할과 BM, 검색일 : 2018년 10월 20일, <http://itmobile.tistory.com/entry/ADNetwork>

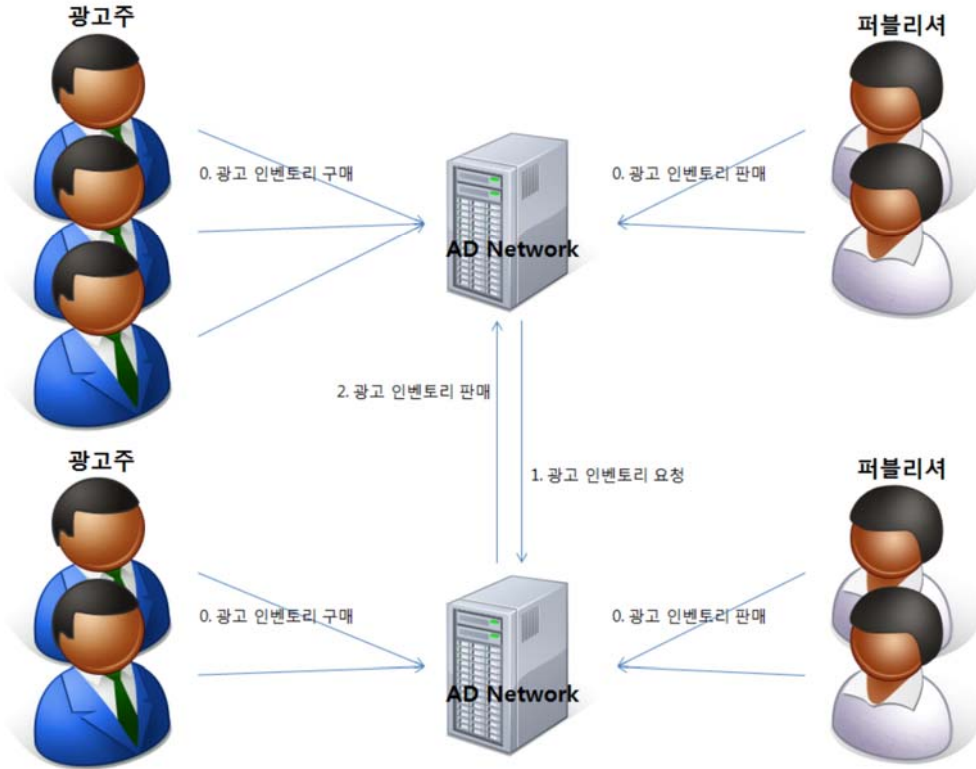
이를 Ad Network라고 부르며, 매체에게 광고주 Ad Server를 연결해주고, 광고주에게는 매체 Ad Server를 연결해서 각 N개만큼의 Ad Server가 필요했다면 이제는 하나의 Ad Network에서 양쪽의 욕구를 충족시키는 플랫폼이 만들어지게 되었다. 이런 Ad Network는 보유하고 있는 광고주나 매체에 따라 서로 다른 특성들을 지니게 되며, 두 거래 주체 사이의 중개인 역할을 통해 수수료를 얻는다.



[그림 2-4] Ad Network 역할

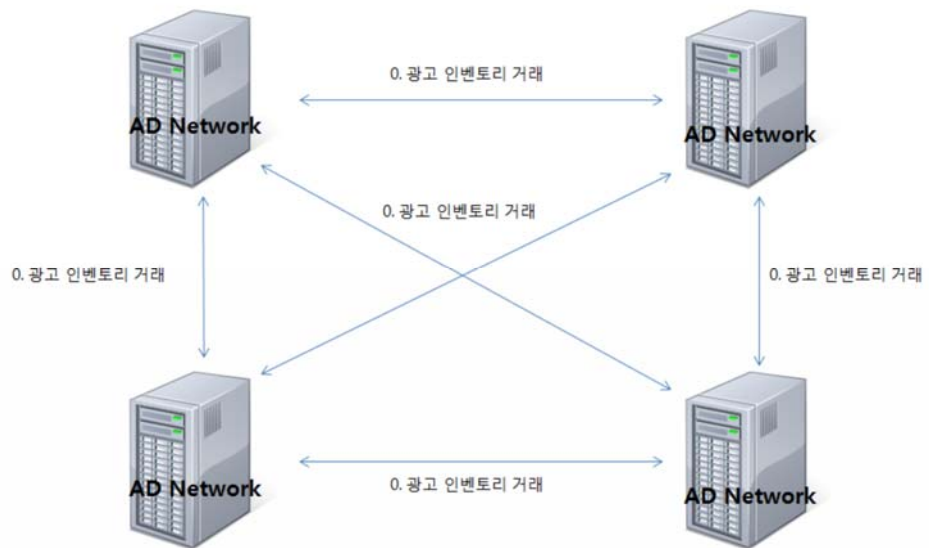
자료: 이현채, (2012년 1월 18일), [Ads. Study] Display Ads, 4편 - AD Network의 역할과 BM, 검색일 : 2018년 10월 20일, <http://itmobile.tistory.com/entry/ADNetwork>

위에서 Ad Network는 보유하고 있는 광고주나 매체에 따라 서로 다른 특성을 지니기 때문에 이용하는 Ad Network에 따라 거래 효율성이 떨어지는 경우가 나타나게 되었다.



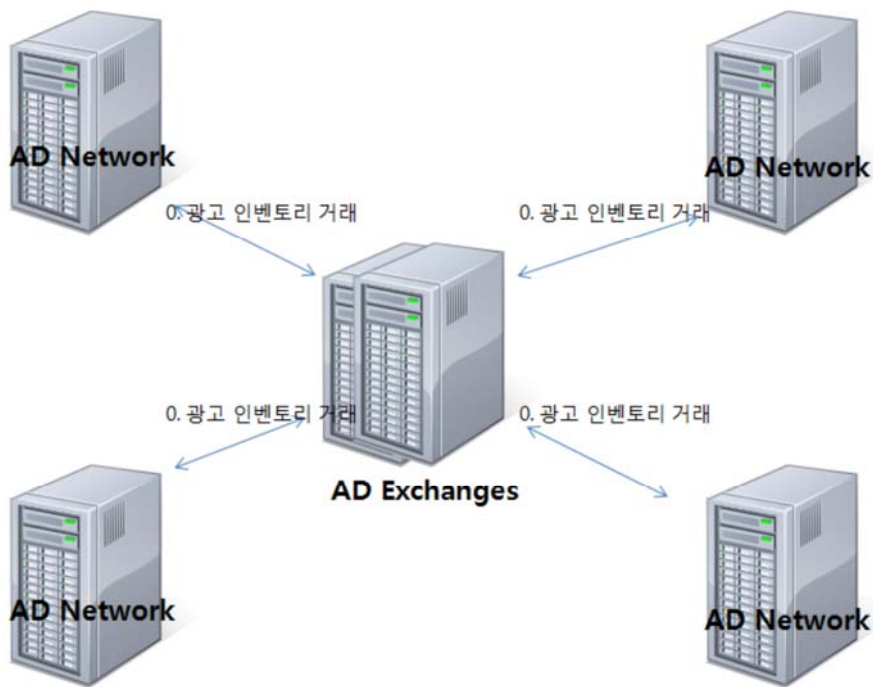
[그림 2-5] Ad Network 사이의 인벤토리 교환

자료: 이현재. (2012년 1월 27일), [Ads. Study] Display Ads. 5편 - AD Exchange의 역할, 검색일 : 2018년 10월 20일, <http://itmobile.tistory.com/entry/ADEXchange>



[그림 2-6] Ad Network 사이의 인벤토리 거래

자료: 이현채. (2012년 1월 27일), [Ads. Study] Display Ads. 5편 - AD Exchange의 역할, 검색일 : 2018년 10월 20일, <http://itmobile.tistory.com/entry/ADEXchange>



[그림 2-7] Ad Exchange 의 등장

자료: 이현채. (2012년 1월 27일), [Ads. Study] Display Ads. 5편 - AD Exchange의 역할, 검색일 : 2018년 10월 20일, <http://itmobile.tistory.com/entry/ADEXchange>

따라서 Ad Network 들의 영역을 거래하는 Ad Exchange가 2005년쯤 생겨나게 되었다. Ad Exchange는 시장에 수요와 공급을 조절하기 위해 다양한 Ad Network를 통합하고, 사용자 방문으로 만들어지는 영역의 광고 송출 기회를 판매하기 위해 경매 방식을 사용한다. 개별 매체와 Ad Network 들은 Ad Exchange에 참여함으로써 이익을 얻을 수 있다. 매체 들은 비-식별 쿠키 데이터를 이용해 사용자를 식별하고, 사용자의 관심사를 파악하여, 사용자 프로파일링을 진행한다. 이렇게 얻어낸 사용자 프로파일링을 그런 사용자를 원하는 광고주들에게 제안함으로써 광고 송출 기회를 판매한다.

매체 입장에서는 영역을 송출 기회를 파는 것이기 때문에 한번 팔리면 특정 영역에 대한 수익은 보장되게 된다. 그러나 광고주 입장에서는 투자수익률을 고려하게 되고, 여러 Ad Exchange와의 거래도 생겨나게 되었다. 특정 영역에 원하는 사용자가 등장했을 경우에 실시간으로 영역을

구매하고자 하는 욕구가 생기게 되었고, 그에 따라 DSP(Demand Side Platform)이 생겨나게 되었다. DSP의 역할은 광고 투자수익률을 높이기 위해, 광고주로부터 받은 정보인 비-식별 쿠키데이터를 바탕으로 각 Ad Exchange에 입찰하는 가격 등을 조정하는 방식으로 광고주의 효율을 높이며 수수료로 운영된다.

광고의 지불 방식이 CPC(Cost Per Click)로 변화함에 따라 클릭이 이루어지지 않으면 광고주도 지불하지 않게 되어, 매체 쪽에서도 이윤극대화를 위한 필요성이 있게 되었다. 그에 따라 SSP(Supply Side Platform)이 생겨나게 되었다. SSP의 역할은 여러 Ad Network로부터 각각의 광고 영역에 대한 CPC 가격을 제안하게 하고, 최고 가격을 선택함으로써 매체의 수익 극대화를 한다.

광고주, DSP, Ad Exchange, SSP, 매체, 사용자 순으로 광고의 흐름이 흘러가며, 온라인 광고 시장은 실시간 입찰 경쟁 시장이 되었다.

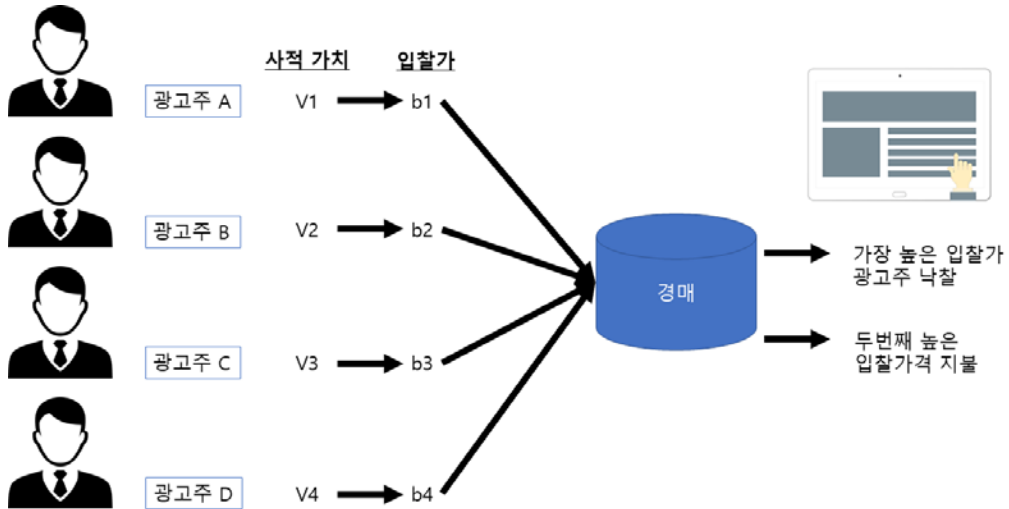
2.2 온라인 광고 시장의 경제학적 특성

온라인 광고 시장은 독점적 경쟁 시장이나 점차 완전 경쟁 시장에 가깝게 이동하고 있다. 완전 경쟁 시장의 조건은 4가지이다. 첫째, 시장 참여자는 모두 완전한 정보력을 갖추고 있다. 둘째, 시장에서 거래되는 재화는 모두 동질적이다. 셋째, 진입과 탈퇴가 자유롭다. 넷째, 충분히 많은 수요자와 공급자가 존재하며 수요자, 공급자 모두 가격을 수용한다. 결국 완전경쟁시장 아래에서 모든 공급자의 시장 지배력은 0이고, 모든 공급자는 가격 수용자가 되며, 모든 재화의 품질이 같기 때문에 오직 가격만이 경쟁력을 가진다.

온라인 광고 시장에서 상품은 각 영역에 대한 광고 송출 기회이다. 이전까지는 1000번 노출 당 비용을 뜻하는 CPM(Cost Per Mille)로 계산을 하였지만, 지불 방식이 점차 CPC(Cost Per Click)로 변화함에 따라 1번 노출 기회에 대한 기대 이익을 추정하는 것이 매우 중요해졌다. 온라인 광고 시장은 인터넷의 발달과 함께 누구나가 광고주가 될 수 있고, 누구나가 매체가 될 수 있다. 초창기에는 매체 파워를 가지고 있는 매체가 유리하였지만, Ad Exchange 등이 등장하고, 지불 방식이 CPC가 됨에 따라 CPC가 유일 가격으로 광고 송출 기회를 결정하게 되었다.

온라인 광고 시장은 수많은 DSP(Demand Side Platform)와 SSP(Supply Side Platform)가 Ad Exchange에서 경쟁에 참여한다. 또한 자신의 선호를 드러내는 Vickery Auction 방식으로 경쟁한다. 그렇기 때문에 특정 경매 순간의 모든 공급자의 시장 지배력은 0이라고 볼 수 있다. 재화의 품질 또한 특정 순간의 광고 송출 기회로 한정 짓는다면 동일하다고 볼 수 있다. 해당 상품의 품질을 가장 높게 측정하는 자가 해당 상품을 갖게 된다.

2.2.1 두 번째 가격 경매



[그림 2-8] 두 번째 가격 경매 예시

Ad Exchange 에서 경매되는 광고 송출 기회는 대부분 두 번째 가격 경매를 채택한다. 이 경매 방식은, 입찰자는 입찰한 가격 대신 2번째로 가장 높게 입찰한 가격을 지불하게 된다.

두 번째로 높은 가격을 지불하는 이유는 Ad Exchange 에서의 상품인 “특정 사용자에게 광고 송출 기회”가 한 번의 게임이 아닌 연속적인 게임이기 때문이다. 만약 첫 번째 가격 경매 방식을 채택하게 되면, 광고주는 자신의 진실 된 가치를 숨기고, 주변 경쟁자의 반응을 살펴 경쟁자보다 약간 우위의 가격을 지불하고자 할 것이다. 자신이 입찰한 가격을 지불하는 첫 번째 가격 경매 방식이라고 하고 예를 들어보겠다. 예를 들어, 2명의 광고주 A, B가 경쟁하는 상황을 가정하자. 각각의 광고주는 자신이 생각하는 특정 광고 송출 기회에 대한 자신에게 내재된 본질적 가치가 있다. 이것을 각각 a와 b라고 하고, a는 8, b는 6으로 가정하자. 그렇게 되면 1번의 게임이 아닌 연속적 게임으로 가정하고 있는 상황에서, 최저 입찰 가격이 2라고 할 때, A 광고주는 입찰 가격이 2부터 시작해서 b 보다 아주 약간 높을 때까지 가격을 올릴 것이고, B 광고주도 역시 2부터 시작해서 b 수준까지 광고비를 계속

올릴 것이다. B 광고주는 b 보다 높은 입찰가를 지불할 수 없기 때문에 그 다음 게임부터 경매에 참여하지 않는다면, A 광고주는 다시 입찰가를 최저 가격인 2로 하여 경매에 참여할 것이고, 그 다음 경쟁부터 다시 B 광고주가 2 보다 아주 약간 높은 가격으로 낙찰을 받은 뒤 가격이 다시 b가 될 때까지 다시 순환된다(Edelman, B. & Ostrovsky, M., 2007).

두 번째 가격 경매 방식이 왜 진실 된 가치를 드러내게 하는지 알아보자(Menezes, F. M. & Monteiro, P. K., 2005). 하나의 광고 송출 기회가 Ad Exchange에서 거래된다고 가정한다. 각각의 광고주 혹은 Ad Network 들은 다른 경쟁자들의 입찰을 관찰할 수 없다. 오직 자신의 입찰가격만을 알 수 있다. 모든 광고주는 입찰가를 자신의 사적 가치(V_i : i번째 광고주의 사적 가치)를 자신이 측정한 광고 송출 기회의 가치 (C_i : i번째 광고주가 해당 광고 송출이 클릭될 것이라 추정한 확률) 로 설정한다고 하자. 모든 광고주에게 클릭의 가치가 동일하다고 가정하고 이를 1로 설정한다. 또한 광고주들이 위험-중립적(Risk-neutral)하다고 가정하자. 그들은 무작위 사건의 기대 가치와 동일 가치를 받는 것에는 무차별하다. 각 광고주가 자신의 사적 가치와 입찰가를 설정하고, 다른 광고주들의 입찰가를 모를 때의 상황에서 기대 이윤을 구해보자.

다른 광고주의 가치는 $[0, +\infty]$ 의 범위를 갖는 임의의 밀도 함수 $f(\cdot)$ 를 갖는 $F(\cdot)$ 에서 추출되었다고 가정하자. 즉, $F_V(v) = P(V \leq v)$ 는 확률 변수 V가 v 보다 작을 확률과 동일하다.

$$\pi_1(v_1, b_i, b(\cdot)) = \begin{cases} v_1 - b_i & \text{if } b_1 > b_i > \max\{b(v_2), \dots, b(v_{i-1}), b(v_{i+1}), \dots, b(v_n)\} \\ 0 & \text{if } b_1 < \max\{b(v_2), \dots, b(v_n)\} \end{cases} \dots\dots\dots [\text{식 2-1}]$$

위의 수식을 보면 첫 번째 광고주가 입찰가를 b_1 으로 설정할 때, 이 b_1 이 다른 상대방의 입찰 가치들 중 최댓값보다 작다면, 기대 이윤은 없고, i 번째 상대방의 입찰가가 가장 높을 때, 그 입찰가보다 커야 사적가치와의 차이가 이윤으로 생긴다.

$$b_1 > b_i > \max\{other\} \text{의 확률은 } \int_0^{b_1} dF(x)^{n-1} = \int_0^{b_1} (n-1)f(x)F(x)^{n-2}dx$$

이다.

$$\pi_1(v_1, b_i, b(\cdot)) = \int_0^{b_1} (v_1 - x)(n-1)f(x)F(x)^{n-2}dx \dots\dots\dots [\text{식 2-2}]$$

위의 기대 이윤에서 첫 번째 광고주가 입찰가를 b_1 으로 설정할 때,
 $b_1 > v_1$ 라면

$$\begin{aligned} \pi_1(v_1, b_i, b(\cdot)) &= \int_0^{v_1} (v_1 - x)(n-1)f(x)F(x)^{n-2}dx \dots\dots\dots [\text{식 2-3}] \\ &+ \int_{v_1}^{b_1} (v_1 - x)(n-1)f(x)F(x)^{n-2}dx \end{aligned}$$

두 번째 항이 음수가 되어, b_1 은 v_1 에 닿을수록 기대 이윤이 증가한다.
반대로, 라면 두 번째 항이 양수가 되어, b_1 은 v_1 에 닿을수록 기대 이윤이
증가한다.

그러므로 $b_1 = v_1$ 일 때, 이윤은 최대가 된다. 그렇기 때문에 상대방의
입찰가를 모르는 상황이라면, 자신의 진실 된 가치와 입찰가를 동일 시 하는
것이 두 번째 가격 경매 방식에서는 지배적인 전략(Dominant Strategy)이
된다.

Ⅲ. 벤치마킹 자료

3.1 자료 소개

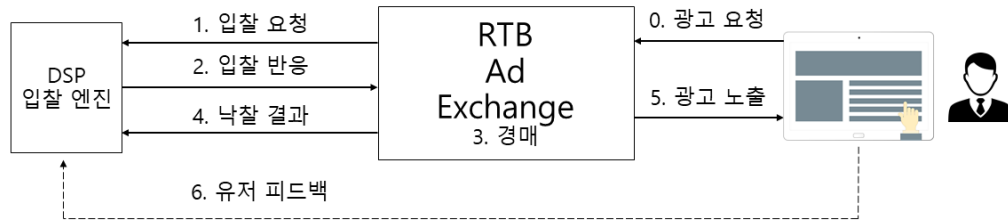
온라인 광고 시장 연구는 연구자들이 데이터에 접근하기 어렵기 때문에 제한적이었다. 그러나 2013년 중국 기업인 iPinYou에서 RTB(Real-Time Bidding) 데이터를 이용한 공모전을 3 시즌 동안 개최하였다. 이 데이터를 공개함으로써 벤치마킹에 사용할 수 있는 자료가 생겼다. Zhang, W., Yuan, S., Wang, J., & Shen, X. (2014)에서는 DSP(Demand Side Platform)에서의 입찰 최적화 문제를 보여주고, 포괄적인 평가 프로토콜 제시한다. 또한 벤치마크 입찰 전략의 실험 결과 제시하였다.

이 데이터는 광고 입찰, 노출, 클릭, 전환을 포함하는 모든 정보를 담고 있다. 광고주별로 캠페인이 하나기 때문에 캠페인별 산업은 1:1이다. 데이터셋 내에는 다양한 산업들이 있다. 광고주 2997을 제외하고 나머지는 CTR(Click Through Rate)이 0.1% 수준이다. 실제로도 광고 산업의 평균은 0.1% 부근이다. 이 말은 천 번을 노출하면 1번 정도 클릭된다는 뜻이다. 광고주 2997은 모바일 광고주라 잘못 놀린 경우도 많이 있다. 전환은 광고주마다 성격이 매우 다르다.

[표 3-1] 광고주 산업

광고주 ID	시즌	산업 카테고리
1458	2	중국의 수직적 e-커머스
2259	3	분유
2261	3	텔레콤
2821	3	신발류
2997	3	모바일 e-커머스 앱 설치
3358	2	소프트웨어
3386	2	국제 e-커머스
3427	2	기름
3476	2	타이어

수직적 e-커머스는 수평적 e-커머스와 달리 여러 카테고리의 상품을 파는 것이 아닌 전문적인 상품은 파는 산업이다.



[그림 3-1] DSP, ADX, 유저 상호작용

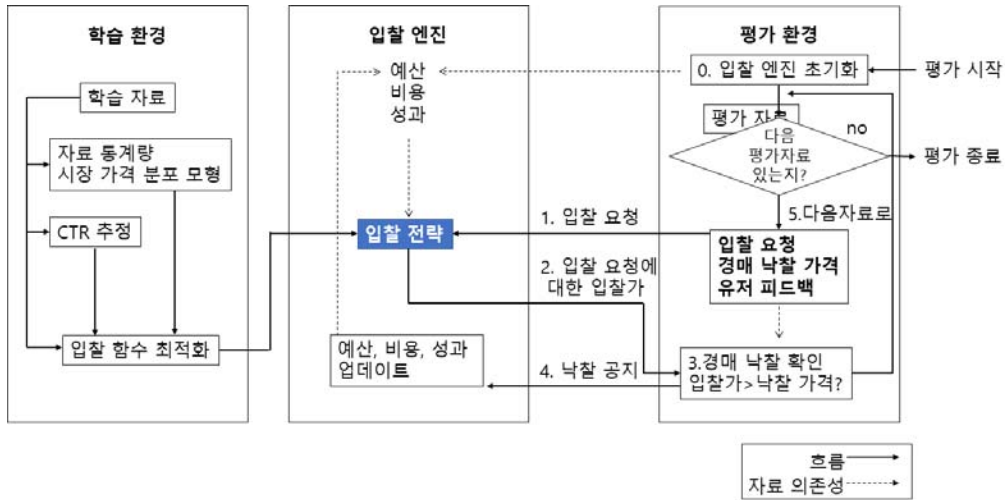
유저가 매체에 등장하면, 매체는 ADX(Ad Exchange)에게 광고 요청을 보낸다. 그럼 ADX는 연결되어있는 DSP(Demand Side Platform)들에게 입찰 요청을 보낸다. 입찰 요청에는 광고 인벤토리의 형태나 유저의 쿠키 등이 포함된다. DSP는 해당 노출 기회에 대한 가치를 평가하고 입찰가격과 제공할 광고를 ADX에게 다시 보낸다. ADX는 DSP들로 모은 입찰 가격을 통해 경매를 진행하고, DSP에게 낙찰 결과를 공지한다. 낙찰된 DSP의 광고가 유저에게 매체를 통해 제공되고, 매체는 해당 노출 기회를 추적하여, 유저 반응(클릭, 전환) 등을 체크해서 DSP에게 제공하게 된다. 이 모든 과정은 100ms 안에 처리된다. iPinYou 데이터 셋은 이 과정을 모방할 수 있도록 되어있다.

유저 프로파일이 2,3 시즌에 들어왔으므로 본 연구에서는 2,3 시즌 자료만을 사용한다. DAAT(Digital Advertising Audience Taxonomy)라는 iPinYou 고유 모델은 데모그래픽, 지오그래픽, 장기 관심사, 시장 구매성향 등을 이용한다. 이 모델은 복잡한 알고리즘을 통해 유저 행동 자료로부터 만들어내어 프로파일을 구성한다. 해당 자료는 안티스팸 필터링이 처리되어 입찰로그에서 의심스런 활동하는 매체를 필터링하였다.

해당 공모전에서의 평가는 각 참가자들은 각 캠페인에 대해 1000 이 주어지고, 입찰 가격이 테스트 데이터 노출 로그의 지불 가격보다 높으면, 지불 가격을 지불하고, 노출했다고 평가한다. 자료에서 모든 가격은

CPM(Cost Per Mille) 당 RMB 위안으로 선형 스케일링되어 원 값을 알 수 없게 처리되었다.

본 연구에서는 Zhang, et al.(2014)에서 제안한 평가 프로토콜을 따른다. 학습 자료를 통해 학습하고, 평가 자료를 통해 평가한다.



[그림 3-2] 평가 프로토콜

[그림 3-2]에서 볼 수 있는 평가 프로토콜은 다음과 같이 진행된다. 각 광고주에 대한 평가 기간의 예산과 입찰 전략을 주어진 것으로 가정한다. 먼저 입찰 엔진을 초기화한다. (0) 예산이나 비용 성과 등을 전부 0으로 설정하는 것이다. (1) 시간 순서에 따라 다음 입찰 요청을 입찰 엔진에 넣는다. 입찰 요청은 페이지에 대한 문맥적 자료와 행동 자료를 포함하고 있다. (2) 예산과 낙찰 가격과 성과에 대한 자료를 갖는 요청에 대해 입찰 전략은 입찰 가격을 계산한다. 만약 비용이 예산보다 높으면 입찰 가격은 0으로 설정하고 남은 모든 입찰 요청들은 건너 뛴다. (3) 노출 로그를 참조하여 경매를 시뮬레이션 한다. 만약 입찰 가격이 기록된 경매 낙찰 가격보다 높고, 최저 가격보다도 높으면, 입찰 엔진은 이 경매에 이긴 것으로 보고, 해당 광고 노출 기회를 얻는다. (4) 경매에서 낙찰되었다면, 이 노출에 대한 로그를 참조하여, 클릭과 전환에 대한 이벤트를 확인한다. 입찰 전략의

성과를 업데이트한다. 비용은 지불 가격이 더해진다. (5) 평가 자료에 남은 입찰 요청이 있는지를 확인하고, 없다면 평가를 종료한다.

오프라인 평가에서 로그를 이용해 재현하는 것의 제한점은 클릭과 전환 이벤트가 낙찰된 경매에 대해서만 가능하다는 것이다. 이것은 잃어버린 경매에 대해서는 어떤 자료도 없음을 의미한다. 그러므로 만약 알고리즘이 입찰가를 높여서 이겼더라면, 성능이 향상되었을 지에 대한 평가를 확인할 수 없다. 다만 이전 연구들에서 오프라인 평가의 편의성을 위해 이 평가 프로토콜이 널리 사용되고 있으므로, 본 연구에서도 관찰되지 않은 유저 피드백은 무시하는 것으로 한다.

iPinYou 가 제공하는 자료 형태는 다음과 같다.

[표 3-2] iPinYou 제공 로그 형태 예시

변수	변수 설명	예시
*1	입찰 ID	93074d812.....1c2374e55a8
2	시간	20131019161502100
3	로그 타입	1
*4	iPinYouID	CAH9FYCtgQf
5	유저에이전트	Mozilla/4.0 (compatible; MSIE 6.0; Windows NT 5.1; SV1; .NET CLR 2.0.50727)
6	IP	119.145.140.
7	지역	216
8	도시	222
*9	ADX	1
*10	도메인	20fc675468.....eda94126da
*11	URL	9c1ecbb8a....36ebf393f7f
12	익명 URL ID	NULL
13	인벤토리 ID	mm_10982364_8930541
14	인벤토리 너비	300
15	인벤토리 높이	250
16	인벤토리 가시성	4번째 노출
17	인벤토리 형태	고정
*18	인벤토리 최저 가격	0
19	소재 ID	7323
*20	입찰 가격	294
*+21	지불 가격	201
*+22	Key Page URL	
*23	광고주 ID	2259
*24	유저 태그	10057,10058,10059

변수 번호에서 *가 붙은 것은 Hash 처리 되어 부호화 되어있다. 이는 익명성을 위해 알아볼 수 없게 비식별 처리한 것이다. 변수 번호에 +가 붙은

것은 입찰 로그에는 존재하지 않으며, 노출/클릭/전환 로그에만 존재하는 변수이다.

첫 번째 변수 입찰 로그는 입찰 요청에 대해서 노출, 클릭, 전환 로그와 연결할 수 있는 값이다. 두 번째 변수는 시간으로 yyyyMMddHHmmssSSS의 1000분의 1초까지 나타내져 있으며, 여기서 시간이나 요일 등의 정보를 추출할 수 있다. 세 번째 변수는 로그 타입이며, 이 노출이 단순노출(1), 클릭(2), 전환(3) 되었는지를 나타내어 최종 모형 성능을 평가할 때 지료를 계산하는 데에 사용된다. 네 번째는 iPinYou에서 관리하는 쿠키 ID이다. 다섯 번째는 브라우저나 기기 등을 알 수 있는 유저 에이전트이다. 모바일로 접속했는지 PC로 접속했는지 등을 구분할 수 있다. 도메인 변수는 광고 인벤토리가 위치해 있는 웹사이트 주소가 된다. URL 변수는 광고 인벤토리 그 자체의 주소가 된다. 익명 URL은 ADX(Ad Exchange)가 URL을 제공하지 않았을 경우에 사용된다. 인벤토리 가시성은 광고 인벤토리가 첫 번째 노출인지 두 번째 노출인지 알 수 없는지를 나타낸다. 인벤토리 형태는 고정, 팝업형태, 배경, 부유(float) 형태인지를 나타낸다. 최저 가격은 광고 인벤토리에 대한 최저 가격으로, 최저 가격 아래로 입찰 할 수 없다. 입찰 가격 변수는 iPinYou가 이 입찰 요청에 대해 입찰한 가격이다. 지불 가격은 경쟁자들로부터 가장 높은 입찰 가격이며, 시장 가격이고 경매 낙찰 가격이다. 만약 입찰 가격이 낙찰 가격보다 높으면, 이 로그는 노출 로그에도 기록된다. 유저 태그 변수는 장기간 관심사를 iPinYou가 해당 유저에 대해 부여해놓은 것이다.

3.2 기초 통계량

[표 3-3] 학습 자료 기초 통계량 1

광고주	기간	입찰수	노출수	클릭수	비용
1458	6월 6일-12일	14,701,496	3,083,056	2,454	212,400
2259	10월 19일-22일	2,987,731	835,556	280	77,754
2261	10월 24일-27일	2,159,708	687,617	207	61,610
2821	10월 21일-23일	5,292,053	1,322,561	843	118,082
2997	10월 23일-26일	1,017,927	312,437	1,386	19,689
3358	6월 6일-12일	3,751,016	1,742,104	1,358	160,943
3386	6월 6일-12일	14,091,931	2,847,802	2,076	219,066
3427	6월 6일-12일	14,032,619	2,593,765	1,926	210,239
3476	6월 6일-12일	6,712,268	1,970,360	1,027	156,088
전체		64,746,749	15,395,258	11,557	1,235,871

각 산업별로 날짜와 입찰 수, 노출 수가 차이가 난다. 그러나 평균적으로 CTR(Click Through Rate) 은 0.1% 에 가깝다. 2997은 모바일 산업이기 때문에 CTR이 0.4%에 가깝다.

[표 3-4] 학습 자료 기초 통계량 2

광고주	낙찰 비율	CTR	CPM	eCPC
1458	20.97%	0.080%	68.89	86.55
2259	27.97%	0.034%	93.06	277.69
2261	31.84%	0.030%	89.60	297.63
2821	24.99%	0.064%	89.28	140.07
2997	30.69%	0.444%	63.02	14.21
3358	46.44%	0.078%	92.38	118.51
3386	20.21%	0.073%	76.92	105.52
3427	18.48%	0.074%	81.06	109.16
3476	29.35%	0.052%	79.22	151.98
	23.78%	0.075%	80.28	106.94

CPM(Cost Per Mille)은 비슷한 전반적으로 비슷한 수준이지만, eCPC(efficient Cost Per Click)는 차이가 많이 난다.

[표 3-5] 평가 자료 기초 통계량 1

광고주	기간	노출수	클릭수	비용
1458	6월 13일-15일	614,638	543	45,216
2259	10월 22일-25일	417,197	131	43,497
2261	10월 27일-28일	343,862	97	28,795
2821	10월 23일-26일	661,964	394	68,257
2997	10월 26일-27일	156,063	533	8,617
3358	6월 13일-15일	300,928	339	34,159
3386	6월 13일-15일	545,421	496	45,715
3427	6월 13일-15일	536,795	395	46,356
3476	6월 13일-15일	523,848	302	43,627
전체		4,100,716	3,230	364,239

평가 자료에는 입찰수가 존재하지 않는다.

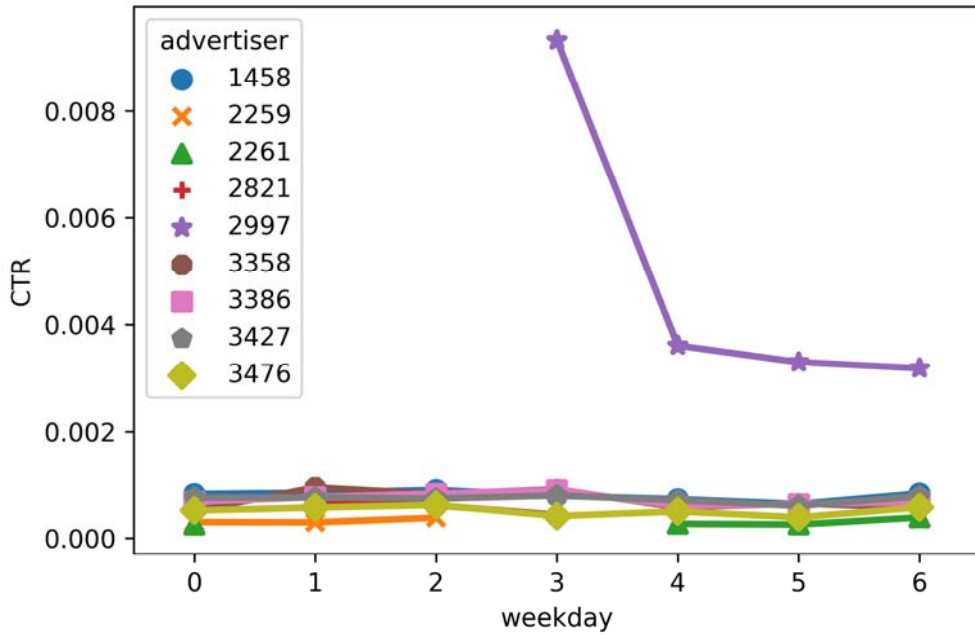
[표 3-6] 평가 자료 기초 통계량 2

광고주	CTR	CPM	eCPC
1458	0.088%	73.57	83.27
2259	0.031%	104.26	332.04
2261	0.028%	83.74	296.86
2821	0.060%	103.11	173.24
2997	0.342%	55.21	16.17
3358	0.113%	113.51	100.76
3386	0.091%	83.82	92.17
3427	0.074%	86.36	117.36
3476	0.058%	83.28	144.46
	0.079%	88.82	112.77

평가 자료에는 입찰수가 존재하지 않기 때문에 낙찰 비율도 없다.

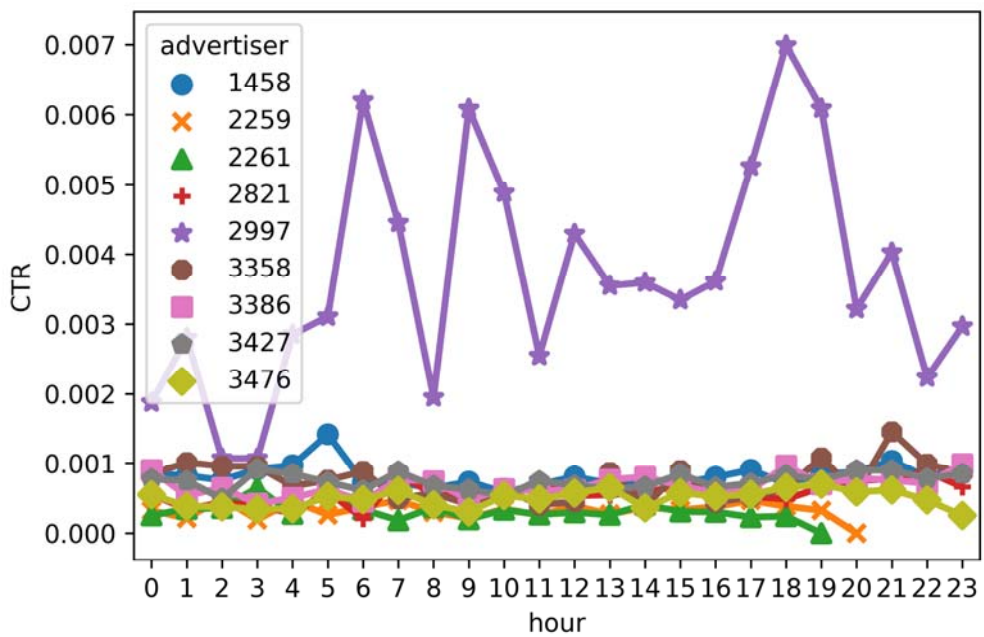
3.3 자료 특성

클릭 확률이 특성 값에 따라 차이가 나는지 본다.



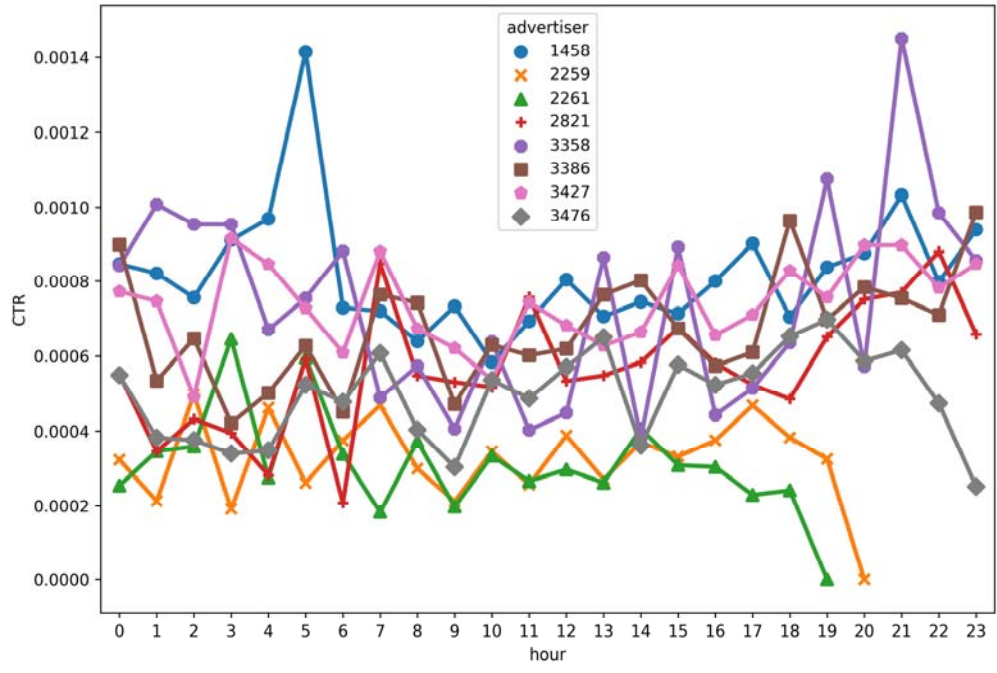
[그림 3-3] 요일별 광고주별 CTR 차이

[그림 3-3]에서 요일별로 광고주별로 CTR(Click Through Rate) 차이가 매우 크게 나타난다. 특히 특정 요일에만 광고하는 광고주들도 존재한다. 0은 일요일을 의미하고, 6은 토요일을 의미한다.



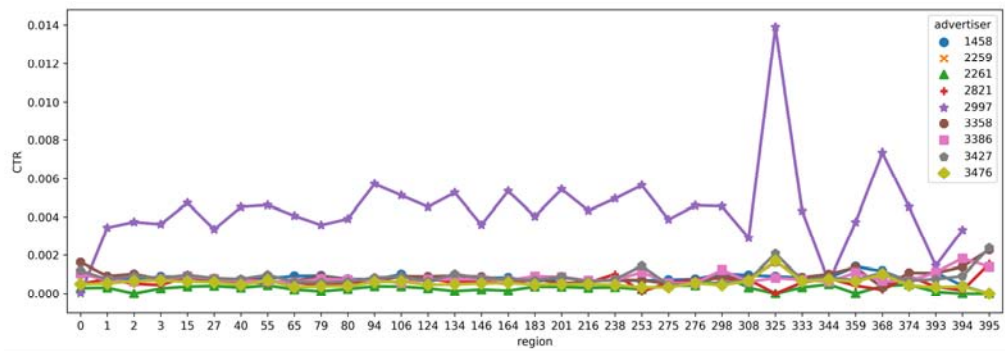
[그림 3-4] 시간대별 광고주별 CTR 차이

[그림 3-4]에서 시간대별로 광고주별로 CTR(Click Through Rate) 차이가 매우 크게 나타난다. 아침시간대에 클릭률이 높은 광고주와 저녁 시간대에 클릭률이 높은 광고주가 보인다. 2997의 모바일 광고가 너무 다른 광고주와 차이가 나기 때문에 빼고 보자.



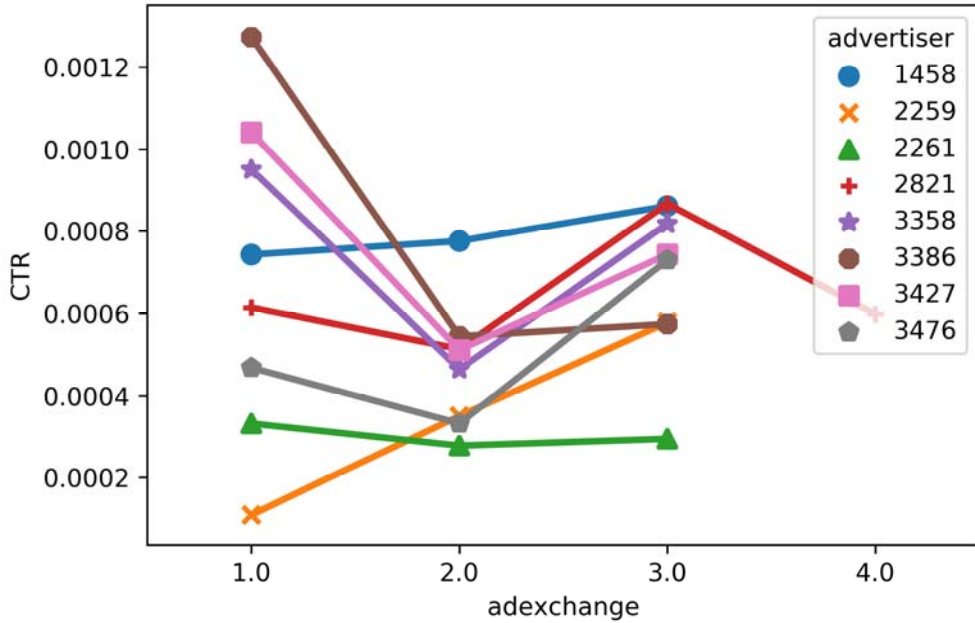
[그림 3-5] 시간대별 광고주별 CTR 차이 (2997 제외)

[그림 3-5]에서 2997의 모바일 광고를 제외하더라도 광고주별로 CTR(Click Through Rate) 차이가 큰 것을 볼 수 있다. 또한 광고주마다 광고를 노출하는 시간대가 다르거나, 미리 소진되어 예산이 부족한 광고주도 있는 것으로 보인다.



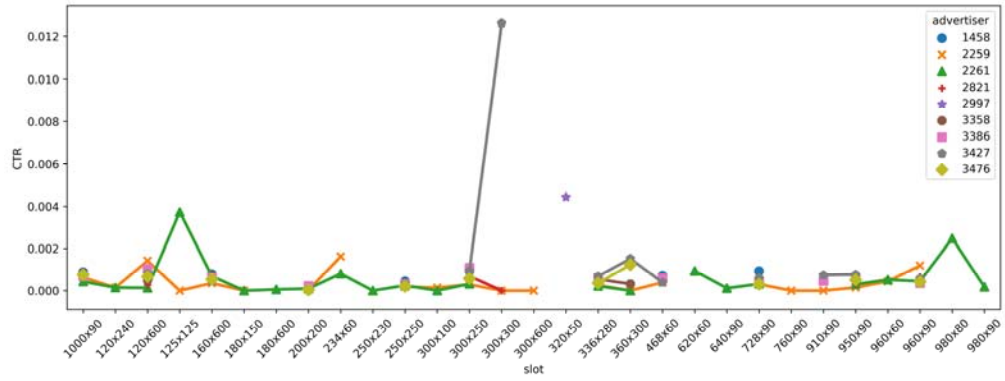
[그림 3-6] 지역별 광고주별 CTR 차이

지역별로도 차이가 난다. 지역별로 변동성이 다르다. 15, 27 같은 지역은 클릭률이 낮은 반면, 325, 368 지역 같은 경우에는 클릭률이 낮고 높음의 차이가 산업별로 크다.



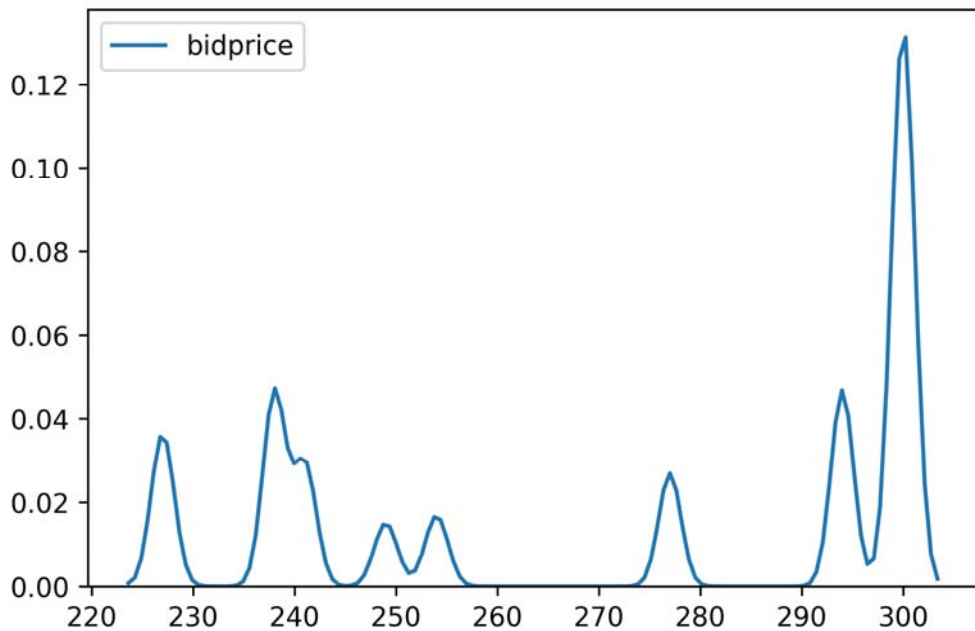
[그림 3-7] ADX별 광고주별 CTR 차이

ADX(Ad Exchange)별로도 차이가 많이 난다. 2259 광고주는 1번 ADX에서는 최하위지만, 3번 ADX에서는 중간 순위이다. 다른 광고주들은 1번 ADX에서 높고 3번 ADX에서 낮음이 보인다. ADX는 그 값 별로 각각 값(플랫폼명, 회사)를 의미한다. 1(Tanx, Alibaba), 2(Adx, Google DoubleClick AdX), 3(Tencent, Tencent), 4(Baidu, Baidu)를 의미한다.



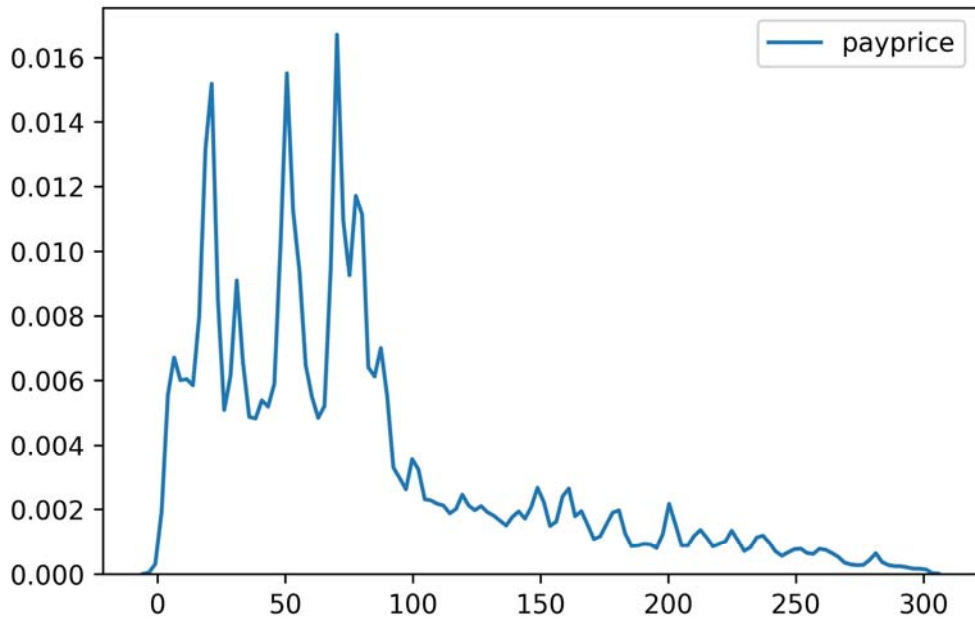
[그림 3-8] 소재 크기별 광고주별 CTR 차이

소재는 일반적인 크기인 300x250이 다른 것들보다 CTR(Click Through Rate)이 높다. 또한 광고주별로 소재의 크기가 매우 다양한 것을 볼 수 있다. 2997 광고주는 모바일 광고주이기 때문에 320x50에만 광고하는 반면, 2261 광고주는 대부분의 소재 크기에 대해 광고를 진행하고 있다.



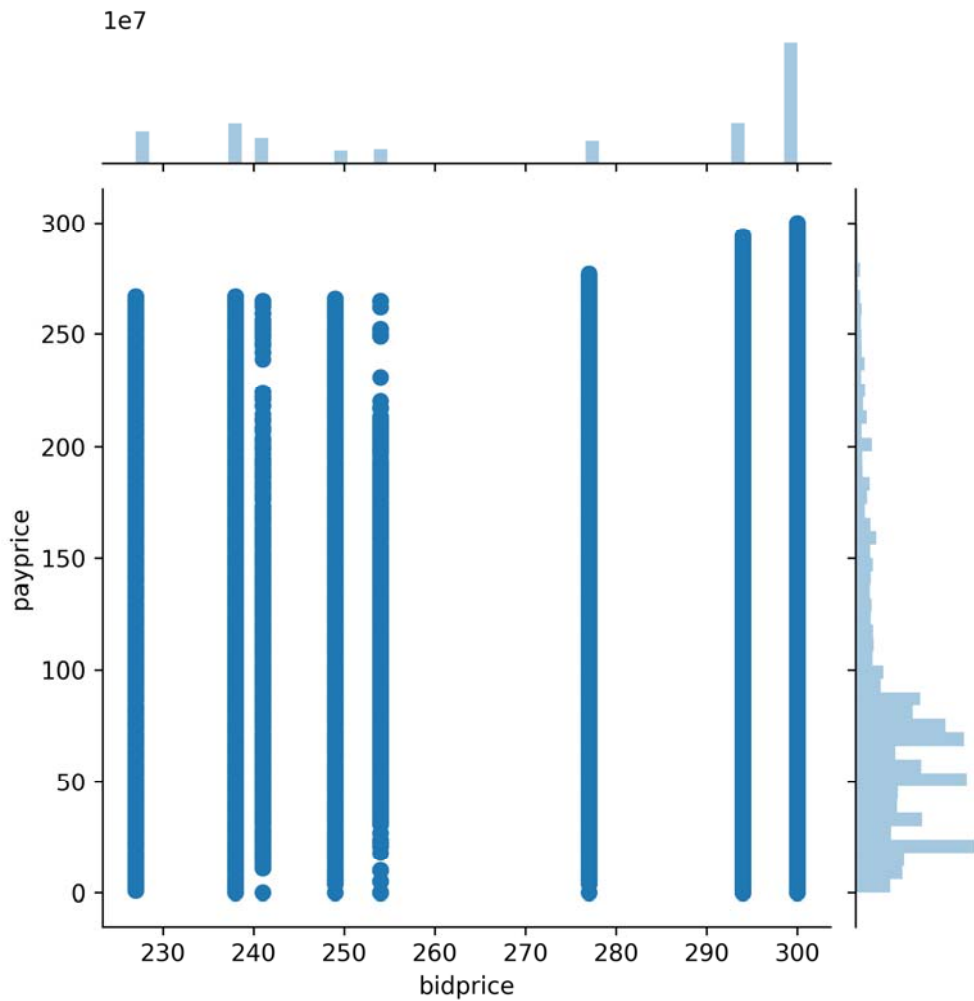
[그림 3-9] 학습 자료에서의 입찰가격 커널 밀도 함수

[그림 3-9]는 학습 자료에서의 입찰가격 커널 밀도 함수를 나타내고 있다. 입찰가격은 전반적으로 넓게 분포하는 것이 아니라 특정 가격대에 밀집되어 있다.



[그림 3-10] 학습 자료에서의 낙찰 가격 커널 밀도 함수

[그림 3-10]은 학습 자료에서의 낙찰 가격 커널 밀도 함수를 나타내고 있다. 낙찰 가격은 0부터 300까지 전반적으로 넓게 분포하고 있어, 입찰 가격과는 차이가 난다. 입찰 가격이 광고주들의 사적 가치를 나타낸다면, 입찰 가격과 낙찰 가격의 차이는 광고주들의 효용이라고 볼 수 있다. 이것을 확인하기 위해 입찰 가격과 낙찰 가격의 결합 분포를 그려본다.



[그림 3-11] 입찰 가격과 낙찰 가격의 결합 분포

결합분포에서 특별한 관계가 보이지 않는다. 결합분포가 단순히 나타내는 것은 입찰 가격에 대해 다양한 낙찰 가격이 존재한다는 것만을 나타낸다. 광고주별로 나눠서 다시 두 변수의 관계를 확인한다.

[표 3-7] 광고주별 입찰 가격의 기초통계량

광고주 코드	평균	표준편차	최소	1사분위	중앙값	3사분위	최대
1458	300	0.00	300	300	300	300	300
2259	288	8.10	277	277	294	294	294
2261	288	8.16	277	277	294	294	294
2821	290	7.07	277	294	294	294	294
2997	277	0.00	277	277	277	277	277
3358	233	6.11	227	227	227	238	241
3386	300	0.00	300	300	300	300	300
3427	236	5.72	227	227	238	241	241
3476	248	6.50	238	238	249	254	254

[표 3-8] 광고주별 낙찰 가격의 기초통계량

광고주 코드	평균	표준편차	최소	1사분위	중앙값	3사분위	최대
1458	69	53.46	0	32	60	80	300
2259	93	74.38	0	29	70	150	294
2261	90	76.89	0	20	63	148	294
2821	89	73.65	0	31	66	147	294
2997	63	60.75	4	18	41	88	277
3358	92	63.92	0	49	77	124	267
3386	77	61.25	0	32	67	90	300
3427	81	57.56	0	43	76	95	267
3476	79	57.95	0	40	73	94	267

노출 로그에 기록된 입찰 가격들은 변동이 매우 적다. 특히 1458 광고주와 2997 광고주 그리고 3386 광고주는 상수 입찰을 하고 있다. 노출 로그에 기록된 입찰 가격들은 단순 입찰 전략으로 보인다.

IV. 클릭 확률 추정

4.1 서론

온라인 광고 시장은 비-식별 쿠키 데이터를 사용하여 사용자 프로파일링을 하고, 이를 바탕으로 개인화 추천을 기반으로 한 광고를 게재한다. 광고주는 CPC(Cost Per Click)로 지불하기 때문에 클릭이 날 만한 사용자에게 대해서만 광고를 내보내고 싶어 한다. 광고주의 입장을 대변하는 DSP(Demand Side Platform)들은 Ad Exchange 내에서 두 번째 가격 경매를 위한 입찰가를 정해야 한다. 입찰가를 정하기 위해서 광고 송출 기회에 대한 가치를 계산해야 하는데, 그 사적 가치는 클릭 확률과 동일하다. 결국 클릭 확률이 광고 송출 기회의 효용이 되는 것이다.

클릭 확률 추정은 주어진 상태 벡터에서 클릭과 같은 반응을 할 사용자의 확률을 추정하는 것이다. 정확한 확률을 얻는 것은 관심사에 관련된 올바른 광고를 사용자에게 제공함으로써 사용자 경험을 향상시킨다. 또한 광고주의 이윤극대화에도 도움이 될 것이다. 그러나 광고에서 클릭 확률 추정은 쉬운 문제가 아니다. 일반적으로 DSP 회사들에서는 이런 사용자 클릭 확률 추정을 위해 관찰치를 매일 수 십 억씩 수집한다. 수 십 억씩 수집된 자료를 학습시키기에 여러 가지 문제점들이 있다. 첫째로는 변수의 수가 너무나 많다. 광고에서 사용되는 자료의 형태에서 대부분의 값들은 명목변수이기 때문에 더미(Dummy) 변수 처리를 해야 하는데 이 수는 기하급수적으로 늘어난다. 둘째, 관찰치가 단순히 많기 때문에 단순 회귀에서처럼 역-행렬을 통해 구하기가 쉽지 않다. 셋째, 메모리 공간과 속도 문제이다. 학습을 하는데만 해도 오랜 시간이 걸리고, 변화가 빠르기 때문에 실시간 학습이 되어야하기 때문에 일괄처리 방식의 학습 방법이 사용될 수가 없다.

클릭 확률은 이항 분류 문제(Binary Classification Problem)이다. 따라서 로지스틱 회귀를 사용하여 문제를 해결 할 것이다. 또한 로지스틱 회귀의 본질적인 문제를 논의한다. 로지스틱 회귀와 관련된 최신 알고리즘을

소개하고, 본 연구의 기여 부분인 GBM(Gradient Boosting Model)의 잎
노드와 K-Means의 변수를 추가한 방법을 논의한다.

4.2 이론적 배경

4.2.1 로지스틱 회귀 (Logistic Regression)

Cox, DR (1958)에서 제안된 로지스틱 회귀는 독립 변수의 선형 결합을 통해 조건부 발생 확률을 추정하는 통계 기법이다. 로지스틱 회귀는 선형 회귀와 마찬가지로 독립 변수와 종속 변수 사이의 선형 관계를 가정한 방정식을 통해, 조건부 평균을 예측하는 것은 동일하다. 하지만 그 가정에서 종속 변수의 연속성이 깨지고, 이항 명목 변수일 경우를 가정하여 문제를 해결하는 방식이다. 두 개 이상의 범주를 가지는 문제가 대상인 경우 다항 로지스틱 회귀 (Multinomial Logistic Regression) 이라 불리며, 순서가 존재하면 순서형 로지스틱 회귀 (Ordinal Logistic Regression)이라고 불린다. 본 논문에서 예측하고자 하는 클릭 확률 예측은 이항 분류 문제이기 때문에 이항 로지스틱 회귀에 대해서 논의하고자 한다.

이항 로지스틱 회귀에서 종속 변수는 선형 회귀와는 달리 종속 변수의 결과가 범위 $[0, 1]$ 로 제한된다. 또한 종속변수의 조건부 확률($P(y|x)$)의 분포가 정규분포가 아닌 이항 분포를 따르게 된다. 대상이 되는 데이터의 종속 변수 y 의 결과가 0, 1만 존재하는 데에 단순 선형 회귀를 적용하면 범위 $[0,1]$ 를 벗어나는 결과를 가져오게 되고, 예측 성능을 떨어뜨린다. 로지스틱 회귀의 핵심은 종속변수의 범위 $[0, 1]$ 를 어떻게 단순 선형 회귀처럼 $[-\infty, +\infty]$ 로 바꾸는 지에 있다. 이를 해결하기 위해 로지스틱 회귀는 $[0, 1]$ 사이의 값을 $[-\infty, +\infty]$ 로 변화시키는 연결 함수(Link Function) $g(x)$ 를 제안하였다. 연결 함수의 형태는 S 자 곡선을 그리는 대부분의 변수가 사용될 수 있으나, 여기서는 가장 널리 사용되는 로지스틱 함수에 대해 서술한다. 로지스틱 함수는 독립 변수가 가질 수 있는 $[-\infty, +\infty]$ 사이의 값을 항상 범위 $[0, 1]$ 사이에 있도록 조정한다. 이것은 오즈비(Odds Ratio)에 로그를 씌움으로서 달성하는데, 이 결과 값을 로짓(Logit)이라고 부른다.

오즈 비란 성공 확률이 실패 확률에 비해 몇 배 더 높은가를 나타내는 값이며 그 식은 다음과 같다.

$$Odds\ Ratio = \frac{p(y=1|x)}{1-p(y=1|x)} \dots\dots\dots [식\ 4-1]$$

분자의 식은 성공 확률을 나타내며, 분모의 식은 실패확률을 나타낸다. 성공확률이 한없이 1에 가까울 0.99999e-20인 경우 값은 무한대에 가까워지고, 성공확률이 한없이 0에 가까울 0.00001e-20 인 경우 값은 0에 가까워진다.

결국 오즈 비를 적용하면 원 함수의 범위 [0, 1]이 [0, +∞] 로 변하게 된다.

이런 오즈비의 결과 값에 로그 함수를 씌우게 되면, 로그 0은 -∞ 이기 때문에 원 함수의 범위 [0,+∞] 가 [-∞, +∞] 의 범위를 갖게 된다.

선형 회귀의 예측 값은 범위가 [-∞, +∞] 이기 때문에 이제 연결 함수를 거치고 나온 값은 선형 회귀처럼 다룰 수 있다.

로지스틱 회귀의 추정 방법은 최대 가능도 방법, 카이 제곱 추정 등이 있으나 여기서는 최대 가능도 방법을 사용하여 회귀 계수를 추정하도록 한다.

최대 가능도 방법이란 임의의 파라미터, 여기서는 회귀 계수 w 가 있을 경우, 주어진 데이터 내에서 발생할 가능성, $P(y|x,w)$ 을 최대로 만드는 방법이다.

회귀 계수 w를 추정하기 위해 확률적 기울기 강하 방법(Stochastic Gradient Descent) 방법을 사용한 추정 방법을 서술한다.

$$\hat{y} = P(y=1|x) = \sigma(w^T x) = \frac{1}{1+e^{-w^T x}} \dots\dots\dots [식\ 4-2]$$

임의의 파라미터 w가 주어졌을 경우, 예측 값 y가 1일 확률은 위의 식과 같다.

$$1 - \hat{y} = P(y=0|x) = \sigma(w^T x) = \frac{e^{-w^T x}}{1 + e^{-w^T x}} \dots\dots\dots [\text{식 4-3}]$$

예측 값 y 가 0 일 확률은 위의 식과 같다. 최적화를 위해 크로스 엔트로피 손실 함수를 적용한다. 크로스 엔트로피 손실 함수는 로지스틱 회귀 모델을 학습하기 위해 일반적으로 사용된다.

$$L(y, \hat{y}) = -y \log \hat{y} - (1 - y) \log(1 - \hat{y}) \dots\dots\dots [\text{식 4-4}]$$

위의 식을 살펴보면 두개의 부분으로 분리 할 수 있고, 다음과 같이 나타낼 수 있다.

$$L(y, \hat{y}) = \begin{cases} -\log \hat{y} & \text{if } y=1 \\ -\log(1 - \hat{y}) & \text{if } y=0 \end{cases} \dots\dots\dots [\text{식 4-5}]$$

위의 식처럼 실제 값 y 가 1일 경우 두 번째 항은 사라지고 $-\log \hat{y}$ 만이 남는다. 만약 예측 값이 1이라면 $\log 1$ 이 되어 손실은 0이 되고, 만약 예측 값이 0에 가깝다면 손실은 $-\log 0$ 이 되어 $+\infty$ 가 된다. 실제 값이 0일 경우 첫 번째 항은 사라지고 $-\log(1 - \hat{y})$ 만이 남는다. 만약 예측 값이 1에 가깝다면 $-\log 0$ 이 되어 손실은 $+\infty$ 가 될 것이고, 만약 예측 값이 0에 가깝다면 $-\log 1$ 이 되어 손실은 0이 된다. 이처럼 크로스 엔트로피는 실제 값과 예측 값이 같을 경우에 손실이 작아지고, 실제 값과 예측 값이 다를 경우에 손실은 기하급수적으로 커진다. 이렇게 구한 손실 함수를 계수 벡터 w 에 대해 미분하여 그라디언트를 구한다.

$$\frac{\partial L(y, \hat{y})}{\partial w} = (y - \hat{y})x \dots\dots\dots [\text{식 4-6}]$$

이 크기가 w 를 아주 조금 변화 시켰을 경우 손실 함수의 변화를

나타내며, w 를 업데이트 한다.

$$w \leftarrow w + \eta(y - \hat{y})x \dots\dots\dots [\text{식 4-7}]$$

여기서 η 는 학습 속도를 조절하는 학습률(Learning Rate) 변수이며, 너무 빠르게 지역 해에 수렴하거나, 너무 느리게 학습되는 것을 조절하게 한다. 일반적으로 클릭 확률 예측에 사용되는 자료는 매우 방대하기 때문에 온라인 학습(Online Learning)으로 이루어지는데, 그 때의 한 번에 변화가 모든 가중치 변화를 가져오지 않도록 조절하는 변수이다. 학습률 변수는 고정된 값이거나, 응용하여 줄어들도록 바꿀 수 있다. 학습의 초반부보다 후반부의 파라미터 미세조정을 위해서 다음과 같이 업데이트 할 수 있다.

$$\eta_t = \frac{\eta_0}{\sqrt{t}} \dots\dots\dots [\text{식 4-8}]$$

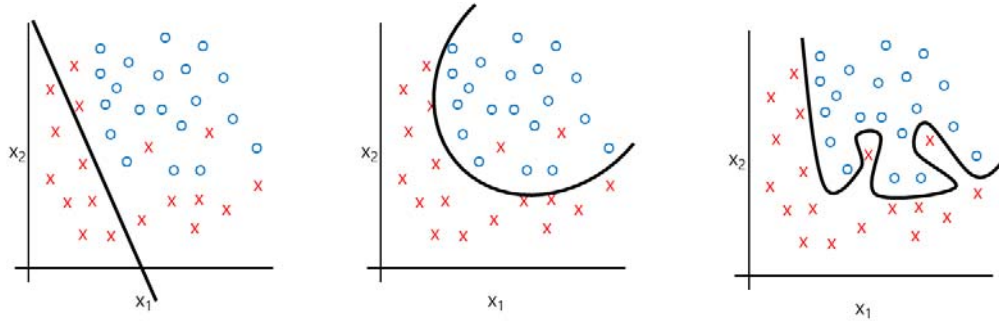
t 번째 반복 마다 학습률의 값을 줄어들도록 만들 수 있다.

4.2.1.1 선형 회귀가 갖는 문제점과 해결책

선형 회귀는 기본적으로 자료 내의 모든 점들의 손실을 최소화하고자하기 때문에, 주어진 자료에 대해 과-적합(Over-fitting)이 발생할 수 있다. 과-적합(Over-fitting)이란 주어진 자료에 대해서만 예측을 잘하고, 새로운 자료에 대해서는 예측 성능이 떨어지는 것을 말한다.

이것은 모델이 데이터의 작은 변화에도 크게 영향을 받는 모델 자체의 분산이 크기 때문에 발생하는 일인데, 이를 해결하는 방법으로써 Tibshirani, R. (1996)에서 제안된 L1 정규화 방법을 소개한다.

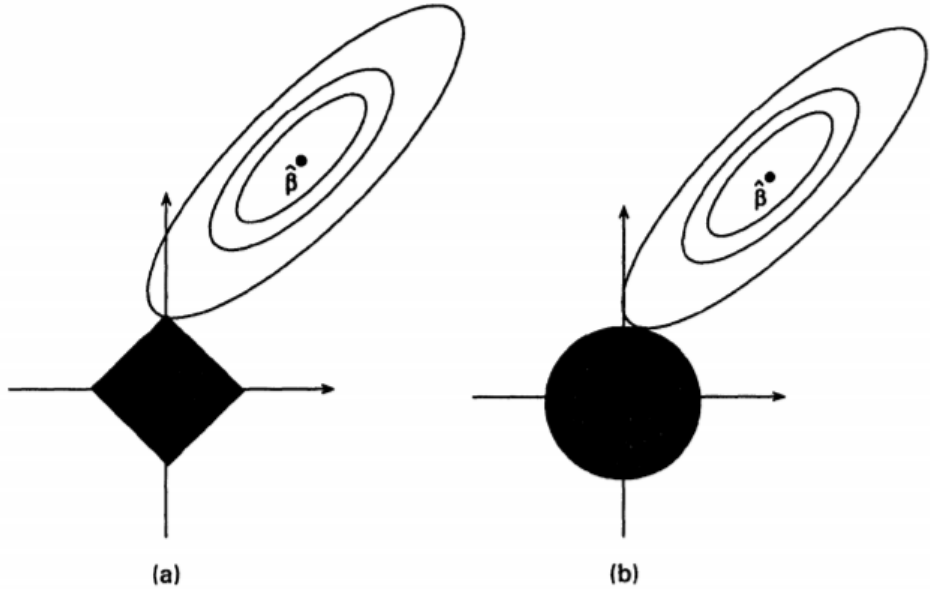
Tibshirani, R. (1996)의 방법은 단순하다. 모델의 분산이 큰 이유는 모델의 복잡도가 높아서 발생하는 일이고, 모델의 복잡도는 파라미터의 개수와 크기 때문에 발생한다.



[그림 4-1] 과-적합 (Over-fitting)

예를 들어 위의 [그림 4-1]을 보면, 가장 왼쪽의 그림은 파라미터의 개수가 너무 적어서 데이터를 잘 분류하지 못하고 있고 손실 함수 값도 클 것이다. 이런 모델은 데이터에 변화가 생겨도 선의 모양이 크게 달라지지 않을 것이다. 가장 오른쪽의 그림은 파라미터의 개수가 너무 많아서 데이터를 완벽히 분리해내고, 손실 함수도 가장 작을 것이다. 그러나 이런 모델은 데이터에 변화가 생기면 선의 모양이 크게 달라질 것이다.

파라미터의 개수가 모델의 복잡도를 결정한다면, 파라미터의 개수가 커지지 않도록 페널티를 주는 방식이 Lasso 방식이다. 이 페널티를 주는 방식에는 어느 파라미터 하나가 너무나 큰 영향을 갖지 않도록 줄여주게 된다. 그 방식으로는 파라미터의 크기를 손실함수에 추가하는 것이다. 가장 많이 쓰이는 두 가지 방식은 L1과 L2 정규화이다.



[그림 4-2] Lasso 와 Ridge 에서 계수 값 추정

자료: Tibshirani, R. (1996) 의 Fig. 2.

L1 정규화는 파라미터에 대해 절댓값 합의 페널티를 주고, L2 정규화는 파라미터에 대해 제곱합의 페널티를 준다. 이 페널티는 손실 함수에 포함되어 파라미터 값이 커질수록 손실이 커지게 만들어, 손실 함수가 자료 내의 손실을 최소화하는 것뿐만 아니라 파라미터 값도 커지지 않게 만들었다. 정규화를 적용한 로지스틱 회귀의 손실 함수는 다음과 같다.

$$L(y, \hat{y}) = -y \log \hat{y} - (1 - y) \log(1 - \hat{y}) + \frac{\lambda}{2} \|w\|_2^2 \dots \dots \dots [\text{식 4-9}]$$

이 손실 함수의 첫 번째와 두 번째 항은 앞서 설명한 것처럼 자료 내의 손실을 줄이고, 새로 추가된 세 번째 항은 파라미터 크기를 줄이게 된다.

새로운 손실 함수로부터 계수 벡터 w 에 대한 그라디언트는 다음과 같이 유도된다.

$$\frac{\partial L(y, \hat{y})}{\partial w} = (y - \hat{y})x + \lambda w \dots\dots\dots [\text{식 4-10}]$$

새로운 그라디언트를 이용해 정규화가 반영된 확률적 기울기 강하는 다음과 같다.

$$w \leftarrow (1 - \eta\lambda)w + \eta(y - \hat{y})x \dots\dots\dots [\text{식 4-11}]$$

4.2.2 선행 정규화를 따르는 로지스틱 회귀 (Logistic regression with follow-the-regularized-leader)

McMahan, H. B., Holt, G., Sculley, D., Young, M., Ebner, D., Grady, J., ... & Chikkerur, S. (2013, August)는 실시간 정규화 방식의 온라인 학습 방법을 제안했다. 온라인 학습에서 관찰치들은 매우 다른 분포에서 올 수 있다. 회귀분석에서 가정하는 I.I.D(Independent and Identically Distributed) 이 지켜지지 않을 수 있다. 이런 경우 단순 SGD(Stochastic Gradient Descent)를 통해 학습을 하게 되면, 매우 빠른 속도로 모델의 파라미터들이 사라질 수 있고, 그렇기 때문에 성능 저하가 나타날 수 있다. McMahan, et al.(2013)에서 제안하는 방법은 이전까지의 그라디언트를 저장하면서, 학습률 또한 그라디언트에 따라 조정되도록 만들어 각 개별 변수들이 자신만의 학습률을 가지고 학습될 수 있도록 바꾸었다. 이 방법의 장점은 처음 등장한 변수에 대해서는 빠르게 학습하고, 많이 등장하는 변수에 대해서는 파라미터 미세 조정이 가능하다는 장점이 있다. 로지스틱 회귀와의 차이점은 모델의 방정식은 동일하나, 최적화 방법이 SGD 에서 위 논문에서 제안된 FTRL(Follow The Regularized Leader)로 변화했다는 점이다. FTRL의 핵심은 다음과 같다. t번째 반복에서 FTRL은 새로운 계수 w를 업데이트 할 때, 다음과 같은 정규화를 진행한다.

$$w_{t+1} = \underset{w}{\operatorname{argmin}} \left(w^T g_{1:t} + \frac{1}{2} \sum_{s=1}^t \sigma_s \|w - w_s\|_2^2 + \lambda_1 \|w\|_1 \right) \dots\dots\dots [\text{식 4-12}]$$

여기서 w 와 곱해지는 g 부분은 t 번째의 누적 그라디언트이고, 는 이전 관찰치의 학습률과의 차이이다. 이처럼 FTRL은 파라미터 업데이트에서 L1 정규화를 같이 적용함으로써, 특정 임계 값을 넘지 못하는 계수 값의 경우에는 0으로 강제 할당하게 된다. FTRL을 사용하면, SGD를 사용할 때보다 0아 아닌 계수값 개수가 줄어들게 되고, 이는 메모리 효율적이 된다.

4.2.3 FM (Factorization Machines)

Rendle, S. (2010)에서 제안된 FM (Factorization Machines)은 고차원 변수들의 상호작용을 고려하기 위해 만들어졌다. 일반적으로 다항 회귀에서 상호작용에 대한 고려는 두 변수간의 곱셈을 새로운 변수로 추가함으로써 이뤄진다.

$$\hat{y} = \sigma \left(w_o + \sum_{i=1}^N w_i x_i + \sum_{i=1}^N \sum_{j=i+1}^N w_{ij} x_i x_j \right) \dots\dots\dots [\text{식 4-13}]$$

위의 식처럼 세 번째 항에서 w_{ij} 가 바로 상호작용을 나타내는 파라미터가 될 것이다. 그러나 여기에서 문제점은 두 변수의 값이 동시에 등장하는 경우에만 저 파라미터가 학습이 될 수가 있다는 점이다. 두 변수가 동시 등장이 없다면, 상호작용은 없는 것으로 학습이 될 것이다. 광고 데이터는 매우 고차원의 특성 벡터를 가지고 있기 때문에 자료의 형태는 매우 희소(Sparse) 해진다. FM에서는 이 문제를 저-차원 상에서 변수 간의 상호작용의 계수를 잠재 벡터에 학습하는 식을 모형 식에 포함하였다. 이를 통해 변수 간 상호작용은 새롭게 만든 상호작용 벡터의 내적으로 구현이 된다. 간단한 2차 상호작용을 포함한 식은 다음과 같다.

$$\hat{y}_{FM}(x) = \sigma \left(w_o + \sum_{i=1}^N w_i x_i + \sum_{i=1}^N \sum_{j=i+1}^N x_i x_j v_i^T v_j \right) \dots\dots\dots [\text{식 4-14}]$$

v_i (i 번째 변수의 저-차원 잠재 벡터) 와 v_j 의 내적은 원래의 상호작용을 모형화 한다. FM 은 행렬 분해(Matrix Factorization)을 통해, 각각의 특성이 저-차원(Low rank)에서 몇 개의 요인으로 축약될 수 있다는 점이다. 이를 통해 잠재 벡터에서 원 행렬을 복원함으로써, 추정되지 않았던 상호작용도 추정이 가능하다는 점이 핵심이다.

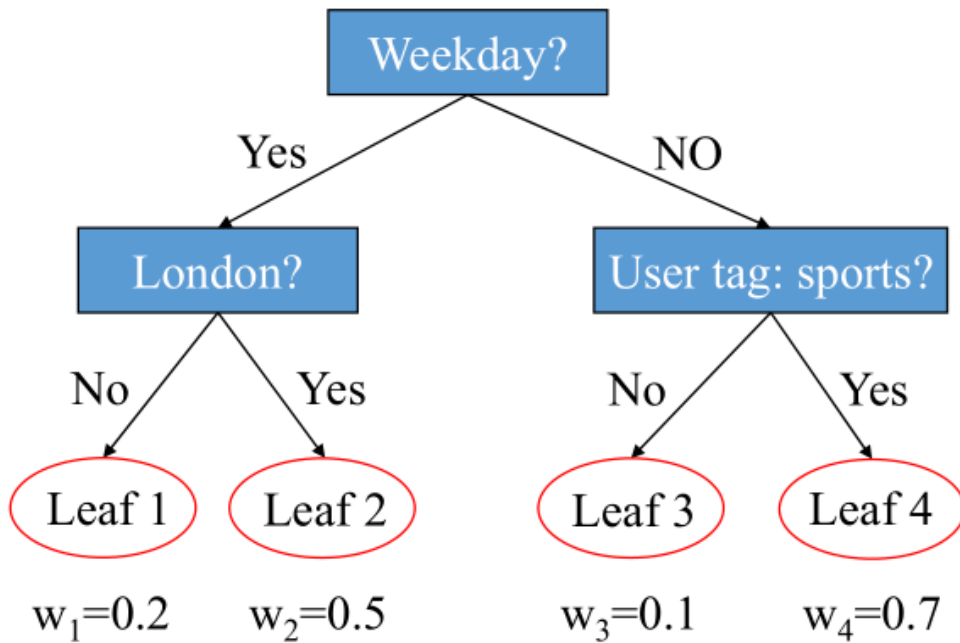
Ta, A.-P. (2015)는 앞서 소개한 FTRL을 FM에 적용하여 CTR 추정 문제에 효과적임을 보였다.

4.2.4 의사결정 나무 (Decision Tree)

의사결정 나무 모델은 이전에 논의한 방법들과 달리 비선형 지도 학습 방법이다. 마찬가지로 CTR 예측에 사용될 수 있다. 주어진 데이터에 모든 변수 모든 임계값에 대해 특정 목적 함수 값을 최적화함으로써 학습을 할 수 있다. 이 특정 목적 함수는 지니 지수(Gini Index)처럼 불순도(Impurity)를 줄이는 방법이 될 수도 있고, 엔트로피(Entropy)처럼 오분류율을 줄이는 방법이 될 수도 있다. 의사결정 나무는 뿌리 노드(Root Node)로부터 출발하여 학습 자료를 각 특성들에 따라 나뉘가며 최종 잎들 중 하나에 도달한다. 잎에 할당된 가중치는 해당 잎의 예측으로써 사용된다. 수학적으로 의사결정 나무는 함수이다.

$$\hat{y}_{DT}(x) = f(x) = w_{I(x)}, \quad I: R^d \rightarrow \{1, 2, \dots, T\} \dots\dots\dots [\text{식 4-15}]$$

$I(x)$ 는 잎을 나타내는 인덱스 함수이며, 관찰치를 잎 t 에 매핑한다. 각 잎은 가중치 w 를 할당받는다. T 는 총 잎의 개수이고, 예측은 매핑된 잎 $I(x)$ 의 가중치로써 할당된다.



[그림 4-3] 의사결정나무의 예

자료: Wang, J., Zhang, W., & Yuan, S. (2016)의 Fig. 4.2.

4.2.5 그라디언트 부스트 나무 모형 (Gradient Boosted Tree Model)

의사결정 나무의 주요 문제는 모형의 높은 분산(High Variance)이다. 모형이 자료에 민감하여, 자료가 조금만 변하더라도 예측값이 높아지는 것을 모형의 분산이 높다고 표현한다. 이러한 불안정성을 줄이고, 과적합을 피하기 위해 실무에서는 여러 개의 의사결정 나무를 결합할 수 있다. 이를 앙상블 기법이라 한다. Friedman, J. H. (February 1999)에서 제안된 그라디언트 부스트 나무 모형은 여러 앙상블 기법 중 부스팅 방법에 해당한다. 부스팅 방법이란 약한 학습기(Weak Learner)를 순차적으로 결합해가면서, 점차 성능이 높은 학습기를 만들어내는 방법이다. 그 중에서 그라디언트 부스트 나무 모형은 잔차를 학습해가며 학습을 하는 방법으로 최근에 널리 사용되는 방법 중 하나이다. 그라디언트 부스트 나무 모형은 이전 약한 분류기에 의해 만들어진 예측치를 지속적으로 조정한다.

$$\hat{y}_{GBTM}(x) = \sum_{k=1}^K f_k(x) \cdots \cdots \cdots [\text{식 4-16}]$$

이 예측치는 K 개의 의사결정 나무 $f_k(x)$ 로부터 나온 예측치의 선형 결합이다. 단순성을 위해, 예측치의 가중치가 동일하다고 가정한다. 그러므로 그라디언트 부스트 나무 모형은 다음의 목적함수를 최소화하기 위한 f_k 들을 찾는 것이다.

$$\sum_{i=1}^N L(y_i, \hat{y}_i) + \sum_{k=1}^K \Omega(f_k) \cdots \cdots \cdots [\text{식 4-17}]$$

목적함수는 실제와 예측의 손실 함수와 각 나무의 복잡성을 조절하는 정규화항 오메가로 구성된다. 공식화하면, 각 단계 $k=1, \dots, M$ 에 대해 다음과 같이 나타낼 수 있다.

$$\begin{aligned} \hat{f}_k &= \operatorname{argmin}_{f_k} \sum_{i=1}^N L(y_i, F_{k-1}(x_i) + f_k(x_i)) + \Omega(f_k) \cdots \cdots \cdots [\text{식 4-18}] \\ F_0(x) &= 0 \\ F_k(x) &= F_{k-1}(x) + \hat{f}_k(x) \end{aligned}$$

이것은 손실 함수의 예측값이 이전 함수들의 결합으로 나타내어지기 때문에 이미 더해지고 학습된 그들의 예측 파라미터를 조정하는 것 없이 새로운 기초 나무를 연속적으로 학습하고 더하는 것에 의한 전진 단계별 추가(Forward Stagewise Additive) 학습이 가능하다. 나무의 복잡도를 조절하는 정규화항 오메가는 다음과 같이 정의된다.

$$\Omega(f_k) = \gamma T + \frac{1}{2} \lambda \sum_{t=1}^T w_t^2 \cdots \cdots \cdots [\text{식 4-19}]$$

람다와 감마는 하이퍼파라미터이다. t는 특성의 인덱스이며, T는 특성의 수이다. 테일러 전개를 통해 예측치 y_k 의 함수로써 목적함수를 고려하여 나타내면 다음과 같다.

$$\begin{aligned} & \sum_{i=1}^N L(y_i, F_{k-1}(x_i) + f_k(x_i)) + \Omega(f_k) \quad \dots\dots\dots [\text{식 4-20}] \\ & \simeq \sum_{i=1}^N (L^{k-1} + \frac{\partial L^{k-1}}{\partial \hat{y}_i^{k-1}} f_k(x_i) + \frac{1}{2} \frac{\partial^2 L^{k-1}}{\partial^2 \hat{y}_i^{k-1}} f_k(x_i)^2) + \Omega(f_k) \\ & \simeq \sum_i^N (\frac{\partial L^{k-1}}{\partial \hat{y}_i^{k-1}} f_k(x_i) + \frac{1}{2} \frac{\partial^2 L^{k-1}}{\partial^2 \hat{y}_i^{k-1}} f_k(x_i)^2) + \Omega(f_k) \end{aligned}$$

여기서 $L^{k-1} = L(y_i, \hat{y}_i^{k-1}) = L(y_i, F_{k-1}(x_i))$ 로 정의하면, 이것은 추가되는 나무 $f_k(x_i)$ 에 독립이므로 제거된다. 목적함수의 기본 학습기 $f_k(x_i)$ 를 의사결정 나무라고 가정하고 다시 표현하면 다음과 같은 식을 얻을 수 있다.

$$\begin{aligned} & \sum_{i=1}^N (\frac{\partial L^{k-1}}{\partial \hat{y}_i^{k-1}} f_k(x_i) + \frac{1}{2} \frac{\partial^2 L^{k-1}}{\partial^2 \hat{y}_i^{k-1}} f_k(x_i)^2) + \Omega(f_k) \quad \dots\dots\dots [\text{식 4-21}] \\ & = \sum_{i=1}^N (\frac{\partial L^{k-1}}{\partial \hat{y}_i^{k-1}} w_I(x_i) + \frac{1}{2} \frac{\partial^2 L^{k-1}}{\partial^2 \hat{y}_i^{k-1}} w_I^2(x_i)) + \gamma T + \frac{1}{2} \lambda \sum_t^T w_t^2 \\ & = \sum_{t=1}^T ((\sum_{i \in I_t} \frac{\partial L^{k-1}}{\partial \hat{y}_i^{k-1}}) w_t + \frac{1}{2} (\sum_{i \in I_t} \frac{\partial^2 L^{k-1}}{\partial^2 \hat{y}_i^{k-1}} + \lambda) w_t^2) + \gamma T \end{aligned}$$

각 의사결정 나무의 잎 노드의 가중치로 변환하면, 해당 노드에 속하는 모든 관찰치는 같은 가중치를 갖기 때문에 잎 노드에 대해 다시 정의할 수 있다. 이것에 대한 최적 가중치를 얻기 위해 잎 노드 t에 대해 0으로 두고 정리하면 다음과 같다.

$$w_t = - \frac{\sum_{i \in I_t} (\partial L^{k-1}) / (\partial \hat{y}_i^{k-1})}{(\sum_{i \in I_t} (\partial^2 L^{k-1}) / (\partial^2 \hat{y}_i^{k-1}) + \lambda)} \quad \dots\dots\dots [\text{식 4-22}]$$

위의 해에 일반적인 최소제곱 오차를 고려해보면 다음과 같다.

$$\begin{aligned} L(y_i, F_{k-1}(x_i) + f_k(x_i)) &= (y_i - F_{k-1}(x_i) - f_k(x_i))^2 \dots\dots\dots [\text{식 4-23}] \\ &= (r_{ik} - f_k(x_i))^2 \\ L^{k-1} &= L(y_i, F_{k-1}(x_i)) = (y_i - F_{k-1}(x_i))^2 = (r_{ik})^2 \end{aligned}$$

r_{ik} 는 단순히 관찰치에 대한 이전 모델의 잔차(Redidual)이다. 그라디언트 부스트 나무 모형은 잔차에 대한 그라디언트로 학습함으로써 해당 이름이 붙게 되었다.

4.2.6 스택킹 앙상블 (Stacking Ensemble)

언급되었던 모든 사용자 반응 예측 모델들은 상호 배타적이지 않고, 결합될 수 있다. He, X., Pan, J., Jin, O., Xu, T., Liu, B., Xu, T., ... & Candela, J. Q. (2014, August)는 로지스틱 모델에 의사결정 나무를 결합하였다. 이 아이디어는 그라디언트 부스트 모델에 의해 예측되는 잎 노드(Leaf Node)를 주어진 입력 특성들의 비선형 변환으로 보고, 이를 새로운 특성으로 로지스틱 회귀에 첨가하는 것이다. 잎 노드들은 데이터를 상호 배타적으로 분리하는 특성이 있으며, 그 각각의 값은 특정 그룹을 나타내고, 그 안에 예측 값은 평균이나 최빈값 등 동일한 분류를 띤다. 마찬가지로 그라디언트 부스팅 나무 모형으로부터 학습된 여러 의사결정 나무들이 있다. 만약 첫 번째 나무가 입찰 요청을 노드 4로 생각하고, 두 번째 나무는 노드 7, 세 번째 나무는 노드 6이면, 특성들 1:4, 2:7, 3:6 이 입찰 요청에 대해 생성되고, 이것을 또 다른 CTR(Click Through Rate) 추정기인 로지스틱 회귀분석에 넣기 위해 해쉬(Hash)된다. 크리테오(Criteo) 캐글(Kaggle) CTR 콘테스트의 후속 솔루션 또한 예측 정확도를 향상시키는 GRBT의 특성들로부터 마지막 예측 값인 잎 노드들을 로지스틱 회귀가 아닌 FFM(Field-aware Factorization Machine)을 사용해 보고했다. Wolpert, D.,(1992)에서 제안된 이렇게 다른 학습기의 예측 값을 또 다른 학습기의

입력 값으로 사용하는 것을 앙상블 기법 중 스택킹(Stacking, Stacked Generalization) 이라 부른다.

4.2.7 K-평균 클러스터링 (K-Means Clustering)

K-평균 알고리즘(K-Means algorithm)은 주어진 데이터를 k개의 클러스터로 묶는 알고리즘이다. 현재 사용되고 있는 표준 알고리즘은 1957년에 스튜어트 로이드(영어: Stuart Lloyd)가 펄스 부호 변조(PCM)를 목적으로 처음으로 고안하였다(Lloyd, S. P., 1982). 무작위로 초기화 시킨 k개의 초기 값을 중심점(Centroid)으로 두고, 각 클러스터와 거리 차이의 분산을 최소화하는 방식으로 학습한다. 이 알고리즘은 비지도 학습의 일종으로, 라벨 없는 입력 값에 라벨을 부여하여 클러스터링을 수행한다.

K-평균 알고리즘은 무작위 초기 값을 중심점에 대해 가장 가까운 클러스터를 할당해놓고, 각 클러스터의 평균점을 재 계산해나가며 수렴할 때까지 학습을 진행한다. 더 이상 모든 점들과 클러스터 내의 평균 중심점이 변동이 없다면 학습을 종료한다. K-평균 방법은 특성을 학습하는 데에도 사용된다. 입력 학습 자료에 대해 K-평균 클러스터링의 K 개수를 달리 해가며 학습하고 난 결과 클러스터링 값을 해당 관찰치의 그룹을 나타내는 변수 값으로 사용함으로써 자료 내의 비선형 관계를 학습한 결과로 생각할 수 있다.

4.2.8 부호화 기법 (Hashing Trick)

Weinberger, K., Dasgupta, A., Langford, J., Smola, A., & Attenberg, J. (2009, June).에서 제안된 해싱 트릭은 주어진 변수 값의 Hash 함수를 적용하고, 이 Hash 함수의 결과 값을 바로 고정된 크기의 변수 배열에서의 인덱스로 사용하는 기법이다. 이 방식은 아무리 많은 변수의 값이라도 충분히 큰 고정된 크기의 배열로 위치시킬 수 있다.

예를 들어, 300만개의 배열을 가정하자. Hash(금요일)의 결과 값이

62352563 이라면, 62352563 / 3000000 을 해서 남은 나머지를 인덱스 값으로 사용하는 방식이다. 이 방식은 충돌이 일어나 서로 다른 값이 같은 배열 인자에 들어갈 가능성이 있지만, 온라인 학습에서는 충분히 많은 학습을 거친 이후에는 해당 배열이 유의미한 파라미터를 학습하는 것으로 알려져 있다.

4.2.9 결측 자료 보정

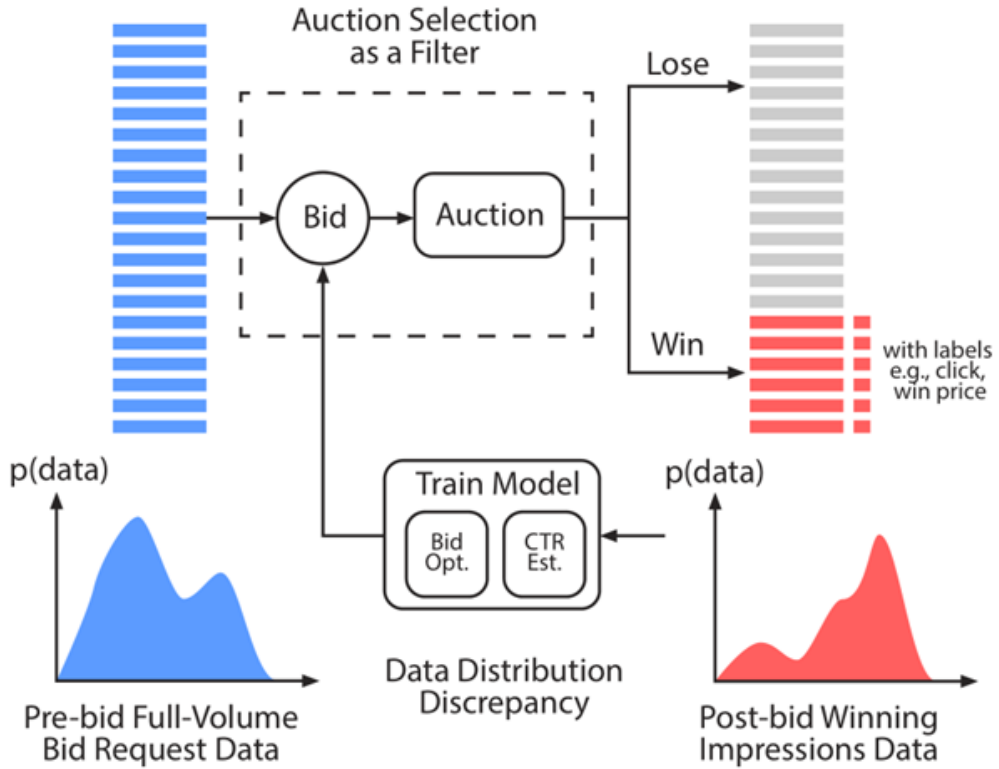
일반적으로 클릭률 추정 학습에 사용되는 학습 자료는 노출 자료이다. 클릭은 입찰 요청이 경매에 낙찰되고 나서 노출이 되어야만, 존재하기 때문이다. 그렇기 때문에 학습 자료는 편향(Bias)되어 있다. 모든 학습 자료는 낙찰된 노출 자료이기 때문이다. 노출자료 분포($q_x(x)$)는 입찰 요청($p_x(x)$)과 낙찰 확률($w(b_x) \equiv P(win|x, b_x) = \int_0^{b_x} p_z^x(z) dz$)을 곱해서 구해질 수 있다.

$$q_x(x) \equiv P(win|x, b_x) \cdot p_x(x) \cdots \cdots \cdots \text{[식 4-24]}$$

위의 방정식은 입찰 전 전체 요청 규모와 노출 데이터의 편향(Bias)를 나타낸다. 학습하고자 하는 예측 모델은 $q_x(x)$ 에서 학습되고, $p_x(x)$ 자료에 대해 예측된다. Zhang, W., Zhou, T., Wang, J., & Xu, J. (2016, August)는 이 편향을 보정하는 방법을 제안했다.

일반적으로 CTR(Click Through Rate) 추정은 다음의 손실 함수를 최소화하는 문제로 공식화된다. 주어진 자료 분포에서 손실을 최소화하는 θ 를 찾아내는 것이 목적이다.

$$\min_{\theta} E_{x \sim p_x(x)} [L(y, \hat{y})] \cdots \cdots \cdots \text{[식 4-25]}$$



[그림 4-4] 학습과 예측 사이의 자료 편향

자료: Zhang, W., Zhou, T., Wang, J., & Xu, J. (2016, August)의 Fig. 1.

Zhang, et al.(2016, August)는 낙찰 함수를 생존 함수로 가정했다. 낙찰 함수를 추정하기 위해, Zhang, et al.(2016, August)에서 제안된 방법을 소개한다. 이를 위해 자료 절단 결측이 없는 상황인 모든 입찰 요청이 낙찰되는 DSP(Demand Side Platform)를 가정한다. 그렇다면 낙찰 확률 $w_0(b_x)$ 는 관측으로부터 직접적으로 얻어진다.

$$w_0(b_x) = \frac{\sum_{(x', y, z) \in D} \delta(z < b_x)}{|D|} \dots\dots\dots [\text{식 4-26}]$$

여기서 z 는 입찰 요청 x' 에 대한 과거 낙찰 가격이고, $\delta(z < b_x)$ 는

$z < b_x$ 일 때 1이고 아니면 0을 반환하는 함수이다. 그렇다면 낙찰 확률은 과거 낙찰 가격이 현재 입찰가격보다 낮았을 경우를 전부 다 더하고 자료 크기로 나누면 된다. 그러나 실제에서는 낙찰이 되지 않는 경우가 더 많다. DSP(Demand Side Platform)가 경매에서 질 경우, 입찰가보다 낙찰 가격이 높았다는 것만을 알려주고, 정확한 값은 모른다. Zhang, et al.(2016, August)는 이 절단된 자료의 편향을 다루기 위해 생존 모형을 사용한다. 생존 모형이란 특정 처리(Treatment) 후에 주어진 시간 동안 환자의 생존율을 예측하기 위한 모형이다. 환자들은 처리 이후에 특정 조사 기간 이후에도 살아남을 수 있기 때문에 최종 생존 기간은 알 수 없지만, 생존 기간이 조사 기간보다 길다는 사실은 알 수 있다. 경매에서 이 생존 모델에 사용되는 낙찰 가격은 환자의 잠재 생존 기간이고, 입찰 가격을 조사 기간으로 간주하면, 유사하게 처리된다. 만약 입찰가가 경매에서 낙찰된다면 낙찰 가격 z 가 관찰된다. 이는 z 일에 환자가 사망한 것을 관측하는 것과 유사하다. 낙찰 되지 않는다면, 낙찰 가격 z 가 b 보다 높았음을 알게 되고, 이것은 환자가 조사를 떠난 것과 유사하다. 이 낙찰 가격 분포 $p_z(z)$ 를 추정하기 위해 캐플란 마이어 누적한계 방법(Kaplan-Meier Product-Limit method)를 사용할 수 있다.

어떤 캠페인이 경매에 N 번 참여했다고 가정하자. 이것의 입찰 로그는 $\langle b_i, w_i, z_i \rangle \quad i = 1, \dots, N$ 처럼 생겼다고 가정하자. 이것은 입찰가격, 낙찰 여부, 낙찰 가격을 나타내는 쌍이다.

자료의 형태를 $b_j < b_{j+1}$ 에 대해서 $\langle b_j, d_j, n_j \rangle \quad j = 1, \dots, M$ 형태로 바꾸면, d_j 는 생존모델에서 b_j 날에 죽은 환자 수와 유사하게 입찰가가 b_{j-1} 일 경우의 낙찰 횟수이고, n_j 는 생존모델에서 b_j 날에 살아있는 환자 수와 유사하게, b_{j-1} 의 입찰가에서 낙찰되지 못한 횟수이다. 이 경우 b_x 의 입찰가를 갖는 광고가 질 확률은 다음과 같다.

$$l(b_x) = \prod_{b_j < b_x} \frac{n_j - d_j}{n_j} \dots\dots\dots [\text{식 4-27}]$$

이는 b_x 날까지의 환자의 생존 확률과 같다. 그러므로 낙찰 확률은 다음과 같다.

$$w(b_x) = 1 - \prod_{b_j < b_x} \frac{n_j - d_j}{n_j} \dots\dots\dots [\text{식 4-28}]$$

b_i	w_i	z_i
2	win	1
3	win	2
2	lose	×
3	win	1
3	lose	×
4	lose	×
4	win	3
1	lose	×

b_j	n_j	d_j	$\frac{n_j - d_j}{n_j}$	$w(b_j)$	$w_o(b_j)$
1	8	0	1	$1 - 1 = 0$	0
2	7	2	$\frac{5}{7}$	$1 - \frac{5}{7} = \frac{2}{7}$	$\frac{2}{4}$
3	4	1	$\frac{3}{4}$	$1 - \frac{5}{7} \frac{3}{4} = \frac{13}{28}$	$\frac{3}{4}$
4	2	1	$\frac{1}{2}$	$1 - \frac{5}{7} \frac{3}{4} \frac{1}{2} = \frac{41}{56}$	$\frac{4}{4}$

[그림 4-5] 낙찰 확률 계산 예시

자료: Zhang, W., Zhou, T., Wang, J., & Xu, J. (2016, August)의 본문 중.

$n_3 = 4$ 의 예제를 보여준다. $z \geq 3-1$ 의 2 케이스, $b \geq 3$ 이다.

$$\begin{aligned}
 E_{x \sim p_x(x)}[L(y, \hat{y})] &= \int_x p_x(x) L(y, \hat{y}) \dots\dots\dots [\text{식 4-29}] \\
 &= \int_x q_x(x) \frac{L(y, \hat{y})}{w(b_x)} dx = E_{x \sim q_x(x)} \left[\frac{L(y, \hat{y})}{w(b_x)} \right] \\
 &= \frac{1}{|D|} \sum_{(x, y, z) \in D} \frac{L(y, \hat{y})}{w(b_x)} = \frac{1}{|D|} \sum_{(x, y, z) \in D} \frac{L(y, \hat{y})}{1 - \prod_{b_j < b_x} \frac{n_j - d_j}{n_j}}
 \end{aligned}$$

로지스틱 회귀에서 각 자료 관찰치에 대한 편의를 제거할 수 있다.

$$f_{\theta}(x) = \frac{1}{1 + e^{-\theta^T x}} \dots\dots\dots [\text{식 4-30}]$$

목적함수를 다음과 같이 다시 정의할 수 있다.

$$\min_{\theta} \frac{1}{|D|} \sum_{(x,y,z) \in D} \frac{-y \log f_{\theta}(x) - (1-y) \log(1 - f_{\theta}(x))}{1 - \prod_{b_j < b_x} \frac{n_j - d_j}{n_j}} \dots\dots\dots [\text{식 4-31}]$$

그렇게 되면, SGD(Stochastic Gradient Descent)에 의해 학습되는 업데이트 규칙도 변한다.

$$\theta \leftarrow \theta + \frac{\eta}{1 - \prod_{b_j < b_x} \frac{n_j - d_j}{n_j}} (y - \frac{1}{1 + e^{-\theta^T x}}) x$$

위의 식으로부터, 낮은 낙찰 가격을 가지면, 학습 집합에서 관찰치를 볼 가능성이 더 낮기 때문에 그래디언트가 커지게 되어, 해당 관찰치에 대해 학습을 많이 한다. 직관적으로, 관찰치 x 가 10% 정도의 낮은 확률로 관찰된다면, 낙찰 확률로 나눠줌으로써 경매에서 지는 잃는 비율 9배를 보정해준다. 만약 극단적으로 낙찰 확률이 높아서 100%라면, 낙찰 분포와 원래 분포 사이의 편향이 없을 것이므로, 그래디언트는 가중되지 않는다.

4.3 연구 설계

4.3.1 데이터와 문제 정의

문제는 주어진 사용자 특성 벡터와 사용자 클릭 여부와의 조건부 확률을 모델링 하는 것이다. 사용자 특성 벡터는 일반적으로 여러 명목 변수로 구성되어 있다.

아래의 표에서 보이듯 대부분의 값은 명목 변수이며, 그 안에 수준(Level)은 매우 다양하게 존재할 수 있다. 요일 변수의 경우 7개의 수준을 갖으며, 시간 변수의 경우 24개의 수준을 갖는다. 이렇게 각각의 명목 변수를 처리하는 일반적인 방법은 더미 변수를 생성하는 것이다. 더미 변수를 생성하는 방법은 해당 수준 값이면 1이고 아니면 0을 갖는 새로운 변수를 만드는 것이다.

[표 4-1] 사용자 특성 벡터의 형태 예시

변수	값
날짜	20180810
시간	16
요일	금요일
IP 주소	255.255.255.255
Ad Exchange	Google
도메인	naver.com
url	http://www.naver.com/
OS	안드로이드
브라우저	Chrome
광고 크기	300x250
광고 ID	12345
기타 등등	

타겟 값은 클릭이 되었으면 1을 갖고 아니면 0을 갖는 벡터이다. 이러한 자료 수집은 Ad Network 에서 하루에도 수 십 억씩 수집되고 있다. 추가적으로 외부 자료를 이용해 데모그래픽 정보를 결합하거나, 비-식별

쿠키 데이터를 바탕으로 프로파일링 한 자료들이 이용된다. 과거의 구매 횟수, 과거의 접속 이력 등이 이용될 수도 있다.

각각의 수많은 명목 변수 아래의 더미 변수를 처리하기 위해서는 3.2.5에서 서술한 Hashing Trick을 사용한다.

클릭 확률 추정에서 가장 문제가 되는 것은 속도 문제이다. 사용자가 매체를 방문하여, 광고 요청을 보낸 뒤 100ms 내에 Ad Network 들은 Ad Exchange에게 광고 입찰가를 보내야하기 때문에 클릭 확률 추정은 100ms 보다 훨씬 더 적은 시간 내에 계산을 끝내야 한다.

4.3.2 분석 방법

본 연구에서 새롭게 제안하고자 하는 것은 Ta.(2015)에서 논의된 FM-FTRL 모델에 추가적인 변수로써 미니 배치 K 평균 클러스터링 알고리즘과 LightGBM 의 잎 노드(Leaf Node)의 결과를 사용하는 것이다. 이것은 4.2.6에서 논의된 스택킹 앙상블 기법의 일종이다. 4.2.5와 4.2.7에서 논의된 것처럼 그라디언트 부스팅 나무 모형의 잎 노드와 K-평균 클러스터링의 결과 값은 주어진 입력 벡터의 비선형 그룹핑으로 볼 수 있다.

학습을 하는 데에 필요한 하이퍼파라미터는 그라디언트 부스팅 나무 모형을 학습하는 데에 학습률 0.05, 나무 개수 50개를 설정하였고, K 평균 클러스터링은 초기 무작위 점을 2^5 개부터 2^9 개까지 늘려가며, 변수를 생성하였다. 나머지 하이퍼 파라미터는 기본 값을 사용했다. FM-FTRL의 경우 알파 베타 값은 각각 1로 설정하였고, FM 가중치들의 대한 알파 베타 값은 0.01과 1로 설정하였다. Hashing Trick에 사용될 전체 가중치 배열은 $2^{22}(4,194,304)$ 개로 희소 행렬(Sparse Matrix)로 학습된다. 사용된 변수는 정수형태의 값을 갖는 10개의 변수를 대상으로 하였다. 모든 모형은 1번의 epoch만을 학습하였다. 학습은 각 광고주별로 학습하였다. 특히 학습할 때 K 평균 클러스터링이나 그라디언트 부스팅 나무 모형을 추가적으로 변수에 삽입하는 경우, 해당 모형들은 전체 학습 자료의 10%를 무작위 추출하여, 학습한 뒤 FM-FTRL의 변수로 추가하였다. 모든 학습 결과는 학습에

사용되지 않은 평가 자료를 대상으로 한 것이다.

4.4 결론

4.4.1 학습 결과

본 연구의 목적은 온라인 광고 시장에서 클릭 확률을 정확히 예측하는 데에 있다. 온라인 광고 시장에서 클릭 확률은 경매에서 사용되는 입찰 요청에 대한 진실 된 가치를 나타낸다. 입찰 요청의 노출 기회에 대한 진실 된 가치를 올바르게 정확하게 추정할수록 기대 수익은 높아지게 된다.

본 연구에서 학습 결과를 평가하기 위해 사용한 지표는 AUROC(Area Under the Receiver Operating Characteristic) Curve 와 RMSE(Root Mean Squared Error) 이다. 또한 주요 클릭 확률 예측 공모전에서 평가 지표로 사용되는 로그손실(LogLoss) 값도 보고한다. AUROC 값은 불균형 분류 문제에서 이상적인 측정 지표로 널리 알려져 있으므로 사용한다.

[표 4-2] FM-FTRL의 학습 결과

광고주코드	시즌	LogLoss	AUROC	RMSE
1458	2	0.0020	0.9644	0.0218
2259	3	0.0028	0.6662	0.0176
2261	3	0.0026	0.5986	0.0168
2821	3	0.0048	0.5940	0.0239
2997	3	0.0226	0.5853	0.0582
3358	2	0.0038	0.9595	0.0267
3386	2	0.0062	0.7660	0.0285
3427	2	0.0029	0.9590	0.0234
3476	2	0.0038	0.9427	0.0233

[표 4-3] FM-FTRL+K-MEANS의 학습 결과

광고주코드	시즌	LogLoss	AUROC	RMSE
1458	2	0.0029	0.8811	0.0214
2259	3	0.0029	0.7464	0.0183
2261	3	0.0027	0.7080	0.0173
2821	3	0.0053	0.6768	0.0252
2997	3	0.0280	0.6809	0.0664
3358	2	0.0041	0.9099	0.0266
3386	2	0.0055	0.7772	0.0269
3427	2	0.0038	0.8713	0.0246
3476	2	0.0038	0.8441	0.0227

[표 4-4] FM-FTRL+GBDT의 학습 결과

광고주코드	시즌	LogLoss	AUROC	RMSE
1458	2	0.0028	0.9269	0.0212
2259	3	0.0029	0.7562	0.0183
2261	3	0.0026	0.7573	0.0173
2821	3	0.0052	0.7211	0.0252
2997	3	0.0270	0.7160	0.0663
3358	2	0.0040	0.9203	0.0265
3386	2	0.0055	0.7857	0.0269
3427	2	0.0037	0.8695	0.0239
3476	2	0.0038	0.8582	0.0227

[표 4-5] FM-FTRL+K-MEANS+GBDT의 학습 결과

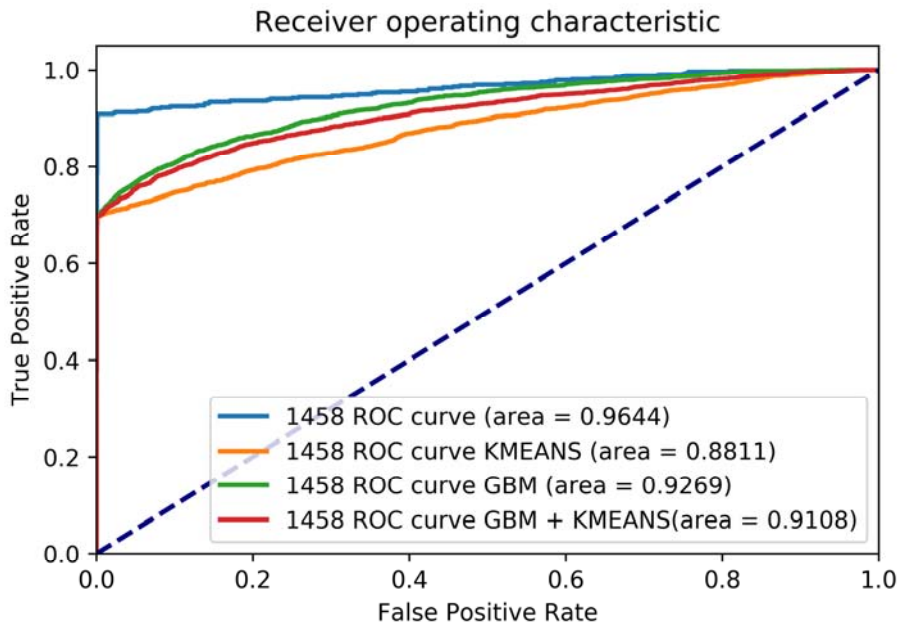
광고주코드	시즌	LogLoss	AUROC	RMSE
1458	2	0.0029	0.9108	0.0213
2259	3	0.0029	0.7792	0.0183
2261	3	0.0026	0.7690	0.0173
2821	3	0.0052	0.7152	0.0252
2997	3	0.0272	0.7195	0.0664
3358	2	0.0041	0.8956	0.0264
3386	2	0.0055	0.7820	0.0269
3427	2	0.0038	0.8783	0.0244
3476	2	0.0037	0.8586	0.0227

학습 결과를 확인해보면, FM-FTRL(Factorization Machine-Follow The Regularization Leader)만 사용했을 때 시즌 2는 평균 0.918 정도의 AUROC(Area Under the Receiver Operating Characteristic)를 갖는 반면, 시즌 3은 평균 0.611 정도를 갖는다. K-MEANS 변수를 추가했을 때 시즌 2는 평균 0.857이고 시즌 3은 평균 0.703 정도이다. GBDT(Gradient Boosting Decision Tree) 변수를 추가 했을 때 시즌 2는 평균 0.872이고, 시즌 3은 평균 0.737이다. 모든 변수를 다 투입하였을 때 시즌 2는 평균 0.865의 AUROC를 갖고, 시즌 3은 0.746 의 AUROC를 갖는다. 이렇게 볼 때, 시즌 2와 시즌 3의 차이가 나는 이유는 Zhang, et al. (2014)에도 보고되어 있듯 시즌 2와 3의 DAAT를 이용한 유저 태그 분류의 차이인 것 같다. 또한 시즌 3의 관찰치가 적기 때문에 FM-FTRL 같은 로지스틱 선형 모형을 사용했을 때 성능이 전반적으로 낮은 것 같다. 시즌 3의 광고주들은 K-MEANS 변수와 GBDT 변수를 포함했을 경우 성능 향상이 눈에 띈다. K-MEANS 변수와 GBDT 변수가 각각 원 자료의 비선형 특성을 전체 자료의 10%만 가지고도 잘 포착하며, 일반화된다고 볼 수 있다. K-MEANS 변수와 GBDT 변수가 차이가 나는 이유는 K-MEANS 는 5개의 그룹 변수만을 추가하였고, GBDT는 50개의 변수를 추가하였기 때문일 수 있다. 광고주 코드별로 LogLoss나 RMSE가 차이나는 것은 산업 특성으로 보인다.

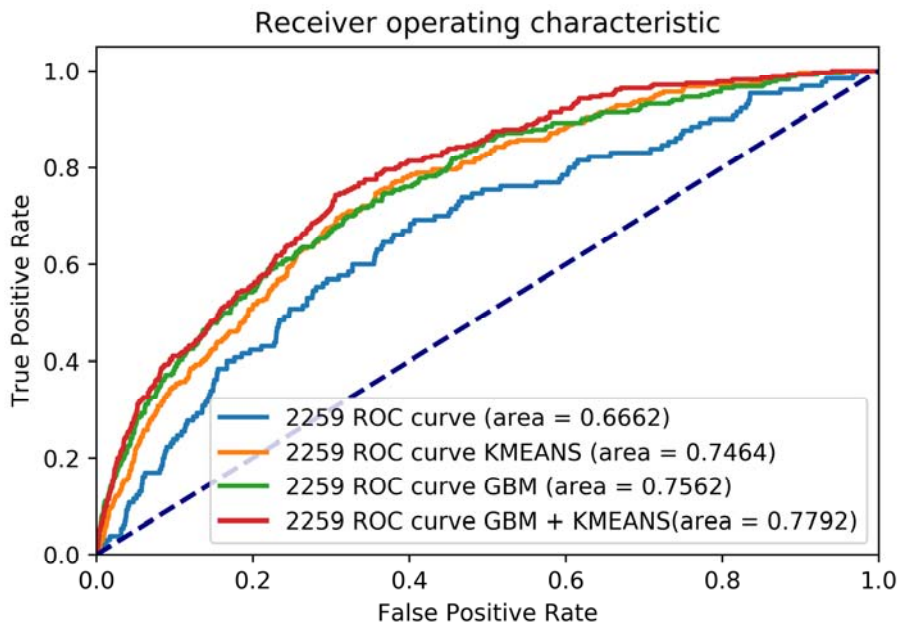
시즌 2에서 K-MEANS 변수와 GBDT 변수가 성능이 안 좋은 이유는 K-MEANS 변수와 GBDT 모두 입력 변수들의 비선형 그룹을 찾는 방식인데, 그 방식이 iPinYou가 분류해놓은 유저 태그와 충돌(Collision)이 일어나는 것으로 추정된다. iPinYou가 자신들의 DAAT를 이용해 유저 태그를 만들 때, 데모그래픽, 지오그래픽, 구매 성향 등의 여러 자료를 사용하는 것이 K-MEANS 같은 방식과 유사한 방식으로 그룹핑 된 것이라면, 본 학습 모델인 FM-FTRL에서 변수 독립이 되지 않아 계수 값을 나눠가짐으로써 성능이 저하된 것일 수도 있다. 다른 이유로는 시즌 2의 관찰치들이 너무 많아서, 10%로는 일반화가 잘 안 되는 것일 수도 있다. 이런 경우 K-MEANS나 GBDT의 잎 노드를 사용하는 것은 적은 관찰치 자료에 더 적합할 수 있다. 이를 나타내어주는 게 시즌 2의 3386 광고주는 시즌 2 자료이면서도 K-MEANS나 GBDT 변수를 추가했을 경우에 성능이 향상되었다.

[표 4-6] 학습 결과 요약표

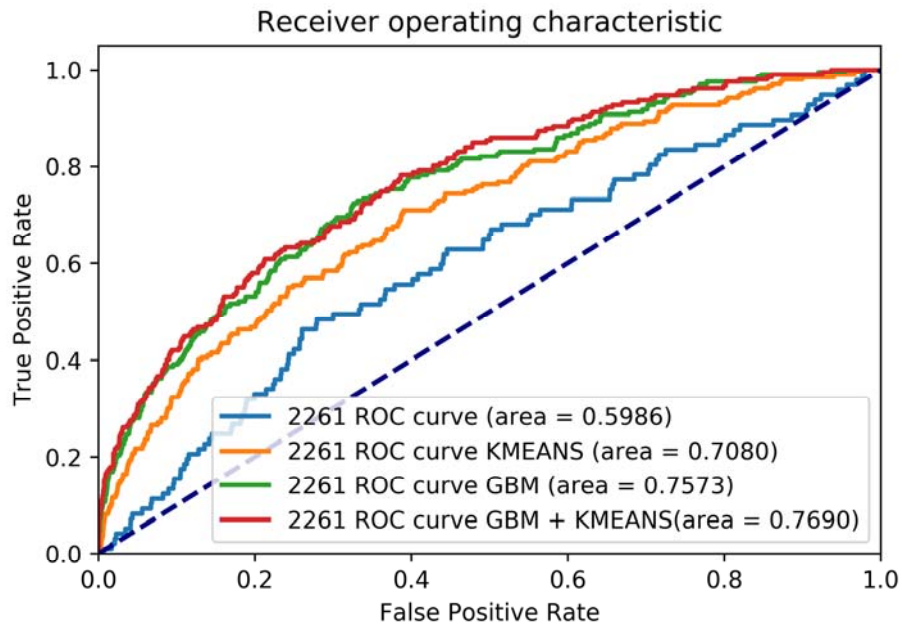
모델	AUROC 전체 평균	향상도(%)
FM-FTRL	0.7817	
FM-FTRL+K-MEANS	0.7884	0.85%
FM-FTRL+GBDT	0.8123	3.92%
FM-FTRL+K-MEANS +GBDT	0.8120	3.87%



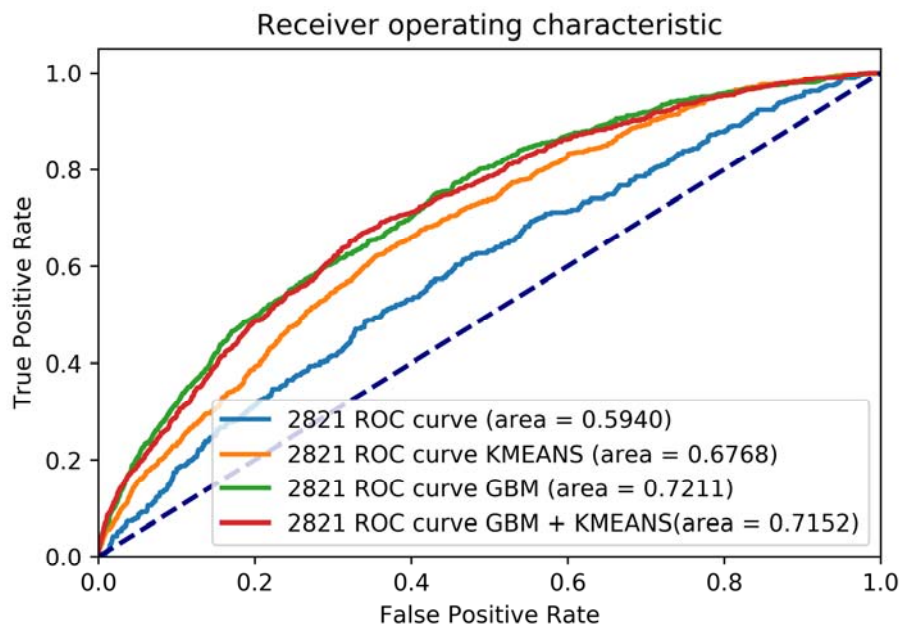
[그림 4-6] 1458 광고주의 ROC Curve



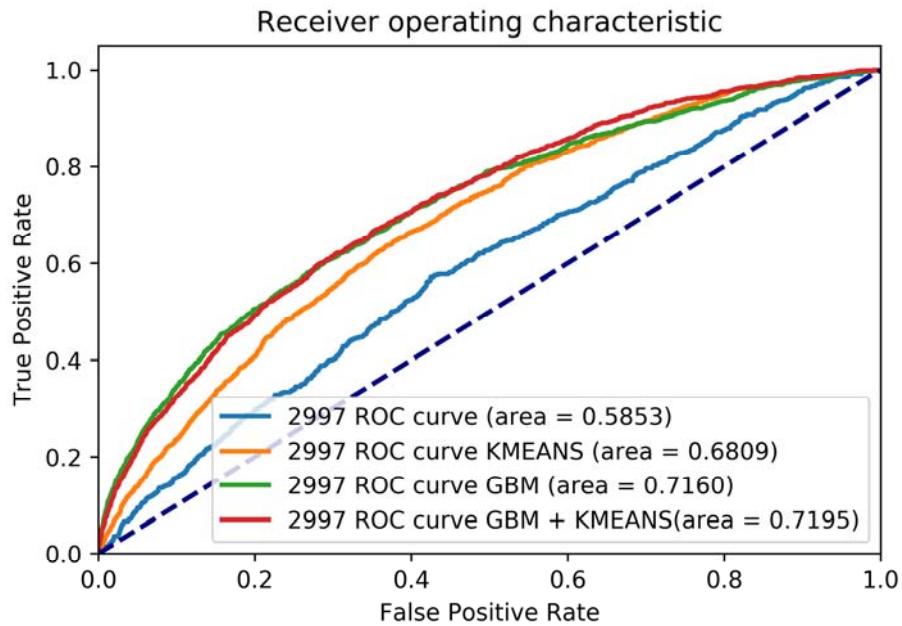
[그림 4-7] 2259 광고주의 ROC Curve



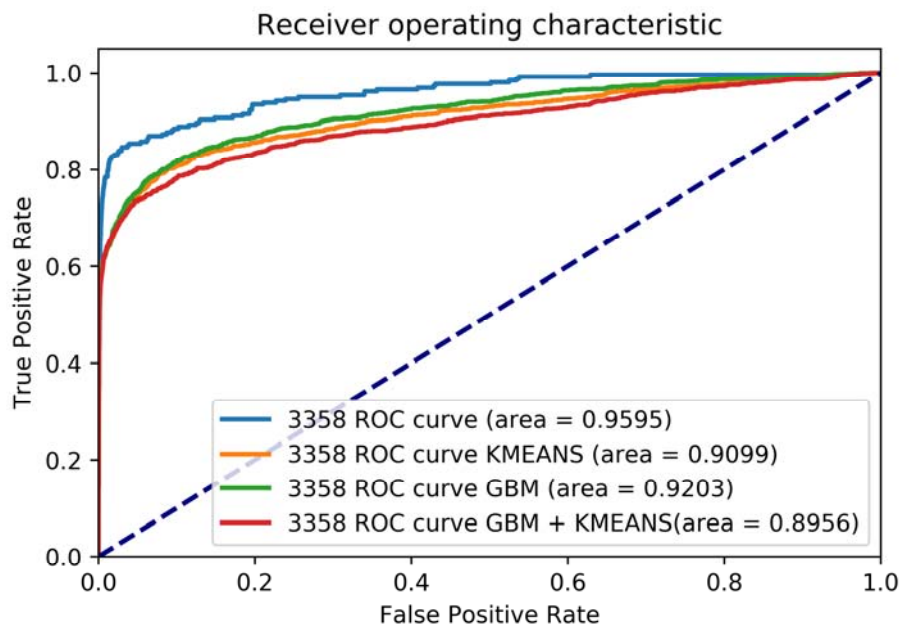
[그림 4-8] 2261 광고주의 ROC Curve



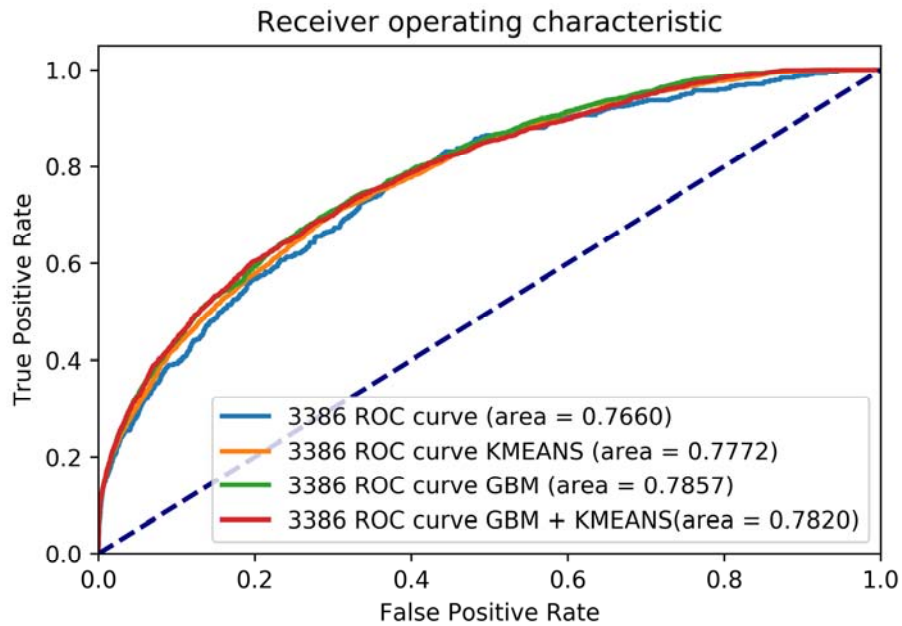
[그림 4-9] 2821 광고주의 ROC Curve



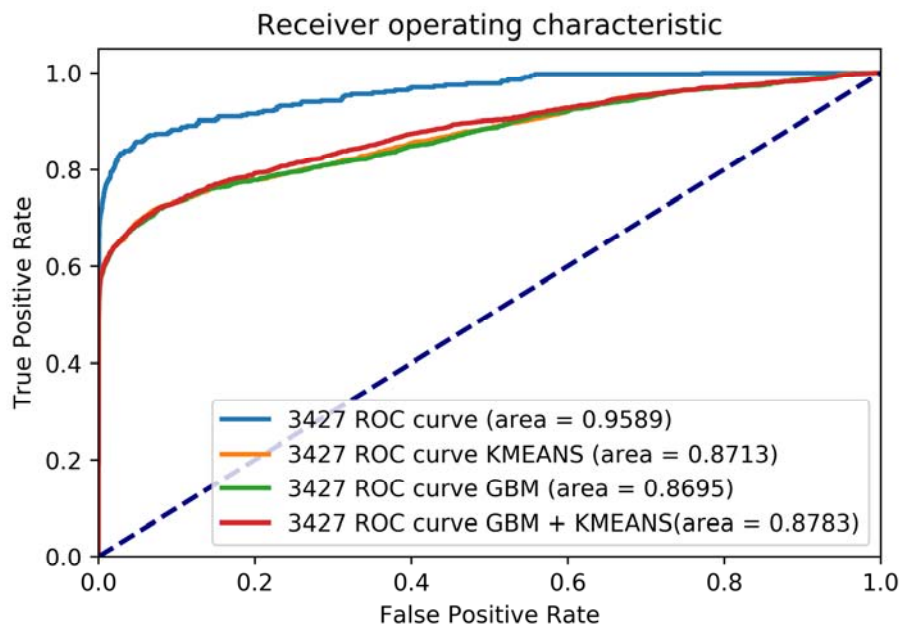
[그림 4-10] 2997 광고주의 ROC Curve



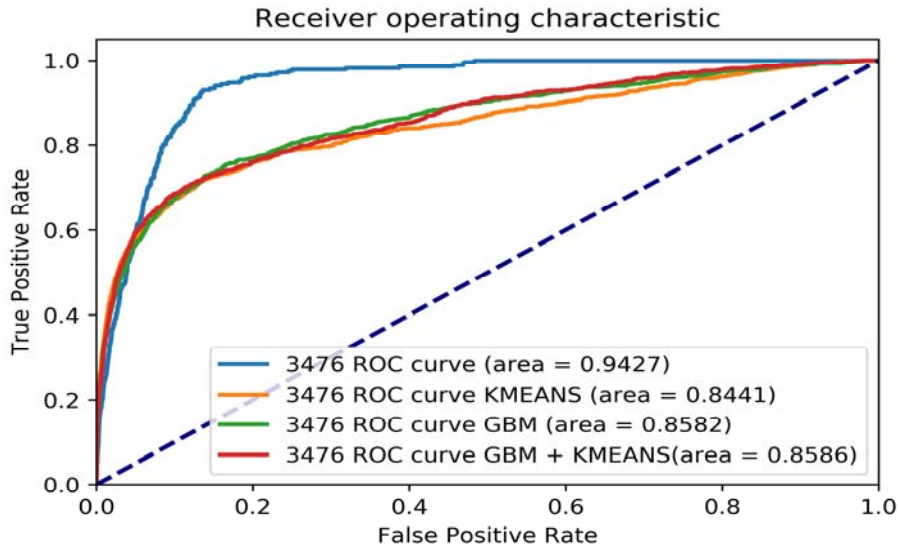
[그림 4-11] 3358 광고주의 ROC Curve



[그림 4-12] 3386 광고주의 ROC Curve



[그림 4-13] 3427 광고주의 ROC Curve



[그림 4-14] 3476 광고주의 ROC Curve

요약하며 4장을 정리하자면, 클릭 확률 추정은 주어진 상태 벡터에서 클릭과 같은 반응을 할 사용자의 확률을 추정하는 것이다. 클릭 확률 추정에는 몇 가지 문제점이 있다. 첫째, 자료로 사용되는 대부분의 변수가 명목 변수이기 때문에 더미 변수 처리를 해야 하는데, 하나의 변수가 갖는 카테고리가 너무나 많기 때문에 더미 변수 처리를 하게 되면 차원이 너무 커진다. 따라서 차원이 기하급수적으로 커지는 문제를 막기 위해, 충분히 큰 고정 차원을 가정하고 해싱 트릭을 이용해 해당 인덱스를 변수로 사용하여 학습시킨다. 둘째, 관찰치가 너무나 많기 때문에 단순 회귀에서처럼 닫힌 해가 나오지 않아 근사 최적화 방법을 사용한다. 또한 광고 산업은 100ms 내에 모든 예측이 끝나야 하는 속도 문제가 훨씬 더 중요하기 때문에 선형 모형을 사용한다. 그에 따라 로지스틱 모형을 소개하고, 로지스틱의 모형의 문제점과 해법, 그리고 개선된 최신 알고리즘인 FM-FTRL 까지 소개하였다. 본 연구의 기여는 널리 사용되는 모델들을 결합함으로써 더 높은 성능을 낼 수 있는 양상블 기법을 소개하고, 벤치마킹 자료를 통해 검증한 데에 있다.

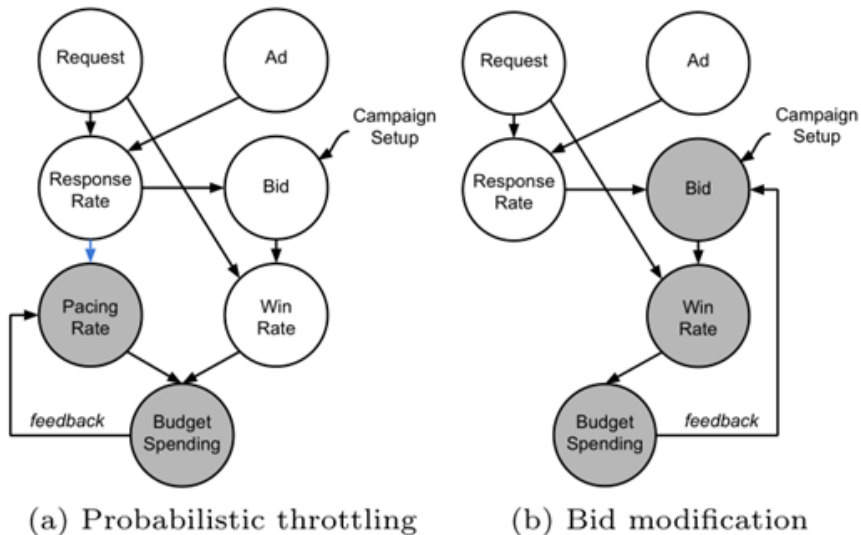
온라인 광고 시장에서 클릭 확률을 정확히 예측하는 것은 기대 수익과 관련이 있다. 잘 예측하고 올바르게 추정할수록 기대 수익이 상승할 것이다.

V. 예산 분배

5.1 서론

광고주는 시간에 따라서 부드럽게 게재되는 예산 할당을 추구한다. 시간에 따라서 부드럽게 예산 분배가 된다면, 같은 광고비를 지속적으로 사용할 수 있고, 더 넓은 잠재고객에게 도달할 수 있다(Lee, K. C., Jalali, A., & Dasdan, A. , 2013, August).

일반적으로 예산 분배 해법은 두 가지 종류가 있다(Xu, J., Lee, K. C., Li, W., Qi, H., & Lu, Q. ,2015, August). 하나는 스로틀링(Throttling) 방법이고, 하나는 입찰 수정(Bid Modification) 방법이다. 스로틀링 기반의 방법은 매 시간 슬롯에서 캠페인이 경매에 참여할 확률을 조정한다. 입찰 수정 기반 방법은 각 경매에 대한 입찰 가격을 적응적으로 조정하며 낙찰 확률을 수정하여 예산 분배 목표를 달성한다.

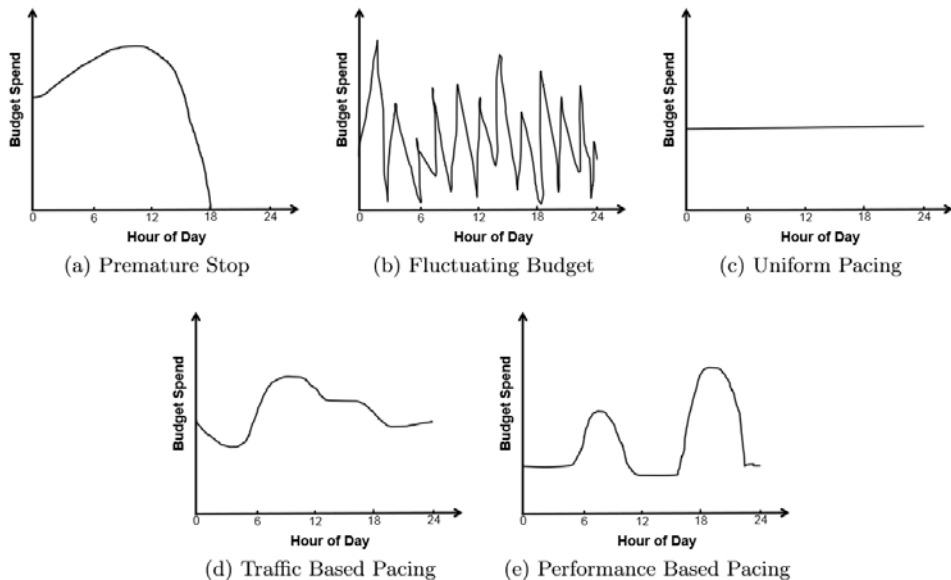


[그림 5-1] 확률적 조정과 입찰 수정 전략

자료: Xu, et al(2015)의 Fig. 2

DSP 에서 광고주의 목적은 다음과 같이 요약될 수 있다. 한 가지는 도달률(Reach Rate)이며, 또 한 가지는 퍼포먼스이다. 브랜딩 광고의 경우, 도달률이 목표이기 때문에 광범위한 오디언스에 도달하는데 예산을 소진하는 것이 목표이며, 그것을 만족하는 동안 비용을 최소로 만드는 것이다. 퍼포먼스 캠페인의 경우, 퍼포먼스 목표(eCPC)를 맞추는 것과 동시에 예산을 가능한 한 사용하는 것이다.

광고주는 자기 자신의 예산 사용 계획을 가지고 있다. 광고주는 여러 매체를 동시에 사용함으로써, 그들의 광고가 광범위한 오디언스에 도달하기를 원한다. 또한 TV나 잡지 같은 전통적인 미디어에서 전달되는 광고와 시너지를 얻기를 원한다.



[그림 5-2] 다양한 예산 분배 형태

자료: Lee, et al(2013)의 Fig. 1

그러나 하나의 캠페인에 대해서도, 광고주의 목적을 맞추는 것은 쉬운 일이 아니다. (a) 와 같은 상황은 조기 소진 되는 것으로 시간이 지남에 따라 예산을 많이 소진하고, 전부 소진한 뒤에는 노출 되지 않는다. (b)는 파동이 있는 형태로 캠페인의 결과를 분석하기에 적절하지 않다. (c)의 형태가

적절해 보이지만, (c)는 두 가지 이슈가 있다. 트래픽 이슈와 퍼포먼스 이슈이다. 트래픽 이슈는 시간에 따라 들어오는 트래픽이 일정하지 않다는 것이다. 그러나 c 같은 유니폼페이싱은 특정 시간대에는 예산을 너무 빨리 소진해버리고, 어느 시간대에는 예산이 부족할 수도 있다. 일반적이고 기본적으로 취하는 방법이지만, 현재는 잘 사용되지 않으며, 트래픽을 반영한 방법은 (d)에 나타나 있다. 두 번째로 퍼포먼스 이슈이다. 잠재 고객에 대한 퍼포먼스가 시간대 별로 일정하지 않다. 퍼포먼스를 높게 내는 시간대에는 예산이 좀 더 많이 분배되어야 하며, 퍼포먼스가 낮은 시간대에는 예산을 줄이는 편이 좋다. DSP(Demand Side Platform)는 개별 증권사와 매우 유사하다. 증권사에서 거래하는 주식은 광고 게재 기회가 되고, 개별 주체가 서로 다른 주식을 거래하며, 기대 이윤을 극대화하기 위해 많은 작업을 한다. DSP를 마치 증권사처럼 생각한다면, 예산 분배는 포트폴리오를 작성하는 것이라 생각할 수 있다. 광범위한 주식 중에서 최적 기대 이윤을 가져다 줄 최적 포트폴리오를 결정하는 문제와 동일시 될 수 있다.

본 연구에서는 예산 분배 알고리즘의 개요를 살펴보고, 최신 알고리즘에 대한 소개와 개선하기 위한 아이디어를 논의한다.

5.2 이론적 배경

5.2.1 부드러운 예산 분배

Lee, et al.(2013)는 예산을 부드럽게 사용하는 방법을 제안했다. Lee, et al.(2013)은 선형 최적화 문제로 예산 분배 문제를 모델링하기 위해 단순한 오프라인 스로틀 기반 해를 제안했다. 캠페인의 일별 예산은 T 개의 슬롯으로 나뉜다고 가정하고, 캠페인의 일별 예산 B 는 각 시점에 걸쳐 할당된다. 만약 입찰 요청 i 가 가치 v_i 와 비용 c_i 를 갖는다면, i 에 대한 입찰을 할지 말지 결정은 x_i 로서 나타낸다. 그러면 선형 최적화 문제는 다음과 같다.

$$\begin{aligned} \max_x \quad & \sum_{i=1}^n v_i x_i \quad \dots\dots\dots \text{[식 5-1]} \\ \text{subject to} \quad & \sum_{j \in I_t} c_j x_j \leq b_t \end{aligned}$$

$\vec{I}_t$ 는 해당 시간 슬롯 T 에 모든 광고 요청의 인덱스 셋이다. 그러나 이 솔루션은 미래의 규모, 가치, 비용의 정보가 없음으로 온라인으로 적용될 수 없다. 그래서 Lee, et al(2013)은 다음과 같은 예산 분배의 온라인 솔루션을 제안했다.

$$\begin{aligned} \min_b \quad & CTR, AR, eCPC \text{ or } eCPA \quad \dots\dots\dots \text{[식 5-2]} \\ \text{subject to} \quad & \left| \sum_t^T s(t) - B \right| \leq \epsilon \quad (\text{total spend}) \\ & |s(t) - b_t| \leq \delta_t \quad (\text{smooth spend}) \\ & eCPM \leq M \quad (\text{max CPM}) \end{aligned}$$

$s(t)$ 는 시간 t 에 실제 지출이다. CPM 안정성 가정에 기반하여, 지출 $s(t)$ 는 다음과 같이 분해될 수 있다.

$$\begin{aligned}
s(t) &\propto \text{imps}(t) && \dots\dots\dots [\text{식 } 5-3] \\
&\propto \text{reqs}(t) \cdot \frac{\text{bids}(t)}{\text{reqs}(t)} \cdot \frac{\text{imps}(t)}{\text{bids}(t)} \\
&\propto \text{reqs}(t) \cdot \text{pacing_rate}(t) \cdot \text{win_rate}(t)
\end{aligned}$$

Reqs(t)는 시간 t에 받은 입찰 요청의 수이다. 페이싱률(t)는 시간 t에 조정될 예산 분배율. 그러므로 예산 b(t+1)을 기대 지출 s(t+1)로 설정하면, 다음 시간 슬롯 t+1의 페이싱률은 다음과 같이 설정된다.

$$\begin{aligned}
&\text{pacing_rate}(t+1) && \dots\dots\dots [\text{식 } 5-4] \\
&= \text{pacing_rate}(t) \cdot \frac{s(t+1)}{s(t)} \cdot \frac{\text{reqs}(t)}{\text{reqs}(t+1)} \cdot \frac{\text{win_rate}(t)}{\text{win_rate}(t+1)} \\
&= \text{pacing_rate}(t) \cdot \frac{b_{t+1}}{s(t)} \cdot \frac{\text{reqs}(t)}{\text{reqs}(t+1)} \cdot \frac{\text{win_rate}(t)}{\text{win_rate}(t+1)}
\end{aligned}$$

실시간 성과와 이전 시점의 페이싱률에 기반하여 계산된다. 요청수와 낙찰 확률의 비율은 과거 정보로부터 가져온다. 이는 이전 시점에 비율을 그대로 사용해도 좋고, 지수 평균을 해서 사용할 수도 있다. 혹은 전일 시점 혹은 일주일의 시점 간 평균 등을 사용할 수도 있다. 이렇게 요청수 비율과 낙찰 확률 비율이 정해지면, 페이싱률은 전적으로 다음 시간 슬롯 T+1 의 예산을 어떻게 정하느냐에 따라서만 달라진다. 시간 슬롯의 예산을 분배하는 방식은 여러 가지가 있다. 예를 들어, 한 가지 전략은 균등 분배이다. 하루 동안 캠페인의 주어진 예산을 균등하게 사용하는 것이다. 이 전략은 쉽게 구현될 수 있다.

$$b_{t+1}^u = (B - \sum_{m=1}^t s(m)) \cdot \frac{L(t+1)}{\sum_{m=t+1}^T L(m)} \dots\dots\dots [\text{식 } 5-5]$$

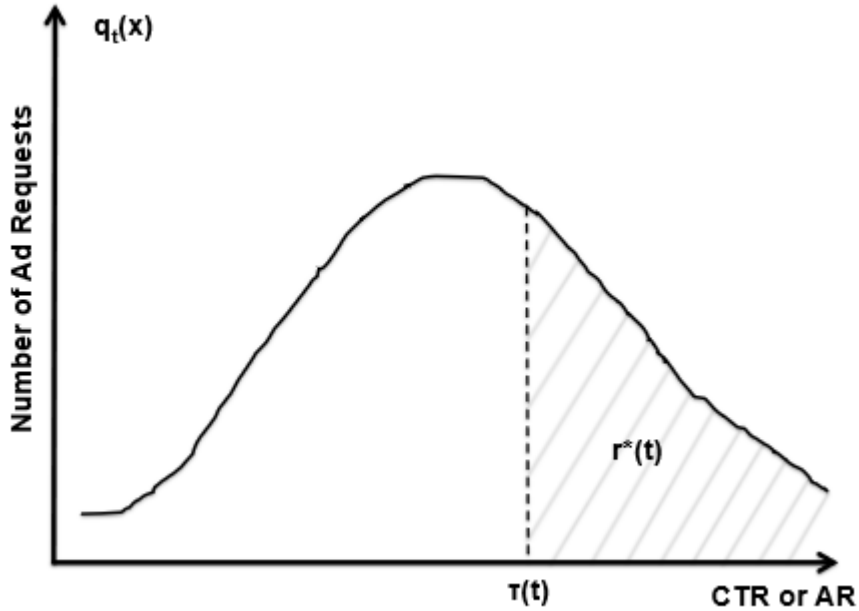
균등 분배 전략은 전체 예산 B에서 현재 시점까지 사용된 예산을 제외한 잔여 예산을 남은 전체 길이와 다음 시간 슬롯 길이의 비율로 곱해줘서 다음 시간 슬롯의 예산을 구한다.

$$b_{t+1}^u = (B - \sum_{m=1}^t s(m)) \frac{1}{T-t} \dots\dots\dots [\text{식 5-6}]$$

남은 시간 슬롯이 매번 동일하다고 가정하면 위와 같이 간단히 쓸 수 있다. 그러나 균등 분배는 캠페인 별 특정 관심 시간대 등의 기회를 포착할 수 없다. 만약 과거 정보를 이용해 캠페인 별 과거 정보의 타임 슬롯 별 기록을 알 수 있다면, 이를 바탕으로 퍼포먼스에 따라 시간대에 더 많은 예산을 부여할 수 있다.

$$b_{t+1}^p = (B - \sum_{m=1}^t s(m)) \frac{p_{t+1} \cdot L(t+1)}{\sum_{m=t+1}^T p_m \cdot L(m)} \dots\dots\dots [\text{식 5-7}]$$

p_t 는 해당 시점의 캠페인의 전환율이나 클릭률 같은 반응 정보이다. 이것은 남은 예산에 시간대별 퍼포먼스의 비율로써 예산이 더 많이 부여될 수 있다. 또한 위에서 본 것처럼 페이싱률은 요청의 비율과 낙찰의 비율이 일정하다고 가정하면, 전적으로 예산 비율에 의해 정해지기 때문에 예산을 높게 잡을수록 페이싱률이 높아진다. 그러나 실제에서는 p_j 의 비율이 0이 되면 예산 소진을 안 하기 때문에, 균등 분배와 섞어 쓰는 전략을 제안한다.



[그림 5-3] 예측 CTR 별 광고 요청 수의 분포

자료: Lee, et al(2013)의 Fig. 2

높은 퀄리티의 광고 요청을 식별하기 위해 저자는 다음과 같은 방법을 제안한다. [그림 5-3]은 예측 CTR(Click Through Rate) 별 광고 요청 수의 분포이다. 만약 모든 광고가 다 클릭된다고 가정했을 때의 최소한도로 필요한 노출 수는 다음과 같다.

$$imps^*(t) = \frac{s(t)}{c^*} \dots\dots\dots [식 5-8]$$

현재 시점에 사용되는 예산을 CPC(Cost Per Click) 단가로 나눈 것이다. 이 최소 노출수를 달성하기 위해 필요한 입찰 수는 다음과 같다.

$$bids^*(t) = \frac{imps^*(t)}{win_rate(t)} \dots\dots\dots [식 5-9]$$

입찰 수를 낙찰 확률로 나눠 최소 입찰 수를 구한다. 유사하게, 저 최소 입찰 기회를 만족시키기 위해서는 다음과 같이 최소 요청수를 구해야한다.

$$reqs^*(t) = \frac{bids^*(t)}{pacing_rate(t)} \dots\dots\dots [식 5-10]$$

입찰 수를 페이싱률로 나눠서 최소 필요 요구 요청수를 구한다. 이렇게 되면, 위에서 구한 $q_t(s)$ 분포에 따라 임계값을 구할 수 있다.

$$\gamma(t) = \operatorname{argmin}_x \left| \int_x^1 q_t(s) ds - reqs^*(t) \right| \dots\dots\dots [식 5-11]$$

CTR별 누적 요청 분포에서 예상 요청수와 차가 가장 적은 지점의 CTR이 된다. 그러나 실무에서는 $q_t(x)$ 가 과거 정보로부터 만들어지기 때문에 지속적으로 변동하는 실제와는 차이가 나게 된다. 따라서 지수 평균을 통해 임계값 파라미터 γ 의 신뢰 구간을 구하여 사용한다.

$$\begin{aligned} \mu_\gamma(t) &= \mu_\gamma(t-1) + \frac{1}{t}(\gamma(t) - \mu_\gamma(t-1)) \dots\dots\dots [식 5-12] \\ \sigma_\gamma^2(t) &= \frac{t-1}{t}\sigma_\gamma^2(t-1) + \frac{1}{t}(\gamma(t) - \mu_\gamma(t-1))(\gamma(t) - \mu_\gamma(t)) \end{aligned}$$

식을 살펴보면 이전 시점 평균 감마에 현재 시점과 이전 사이의 평균만큼 더해 새로운 평균을 구하고, 분산도 마찬가지로 이전 시점에 증가율에 차이의 평균만큼 더해 새로운 분산을 구한다. 이렇게 구한 값을 이용해 요청에 대한 입찰확률을 결정한다. 예를 들어, 처음 광고 요청이 들어오면 이 요청에 대한 예측 CTR(Click Through Rate) 값을 구하고 이 값이 감마의 상위 경계보다 높으면 바로 입찰하고, 그렇지 않으면 페이싱률에 따라 입찰에 참가하게 된다.

5.2.2 스마트 예산 분배

Xu, J., at el.(2015, August)에서 제안된 것 같은 스로틀 기반 솔루션은 온라인 페이싱률 제어와 최적화의 프레임워크에서 일반적이다. 입찰 수정 기반 방법들은 조정된 입찰 가격들이 키워드 수준에 설정된 스폰서 광고에서 연구되었다[Mehta, A., Saberi, A., Vazirani, U., & Vazirani, V. (2007), Borgs, C., Chayes, J., Immorlica, N., Jain, K., Etesami, O., & Mahdian, M. (2007, May)]. RTB 디스플레이 광고에서, 피드백 통제 기반 방법들은 예산 페이싱 작업에서 입찰 수정을 위해 채택되었다[Chen, Y., Berkhin, P., Anderson, B., & Devanur, N. R. (2011, August), Karlsson, N., & Zhang, J. (2013, June), Zhang, W., Rong, Y., Wang, J., Zhu, T., & Wang, X. (2016, February)].

Xu, J., at el.(2015, August)는 확률적 스로틀 해법을 휴리스틱 방법을 사용해 개선하였다. 각 시간 슬롯에 대해 부여 되어 있는 예산이 있고, 소비된 비용이 있다면, 이 비용과 예산의 차이를 줄이는 최적화 문제를 설정하였다. 해 공간(Solution space) 를 줄이기 위해, 반응 예측 모델을 만들어, 비슷한 반응 보이는 캠페인을 하나의 그룹으로 보고, 페이싱률을 동일하게 적용하였다. 두 번째는 실시간 피드백 자료를 이용해 최적 해에 근사하기 위해 동태적으로 그룹 페이싱률을 조정하였다.

각 캠페인에 대해 DSP는 적합한 광고 요청을 받으면, 광고 요청이 어느 광고 요청 그룹에 속하는 지 결정하고 해당 레이어를 참조하여 페이싱률을 얻는다. DSP는 입찰 최적화 모듈에 의해 주어진 가격에서 페이싱률과 동일한 확률로 광고 요청을 입찰한다.

$$\Omega(C,B)=\sqrt{\frac{1}{K}\sum_{t=1}^K(C^{(t)}-B^{(t)})^2}.....[\text{식 5-13}]$$

$$C=\sum_i=s_i\times w_i\times c_i.....[\text{식 5-14}]$$

$$B=(B^{(1)},...,B^{(K)}).....[\text{식 5-15}]$$

$$P = C / \sum_i s_i \times w_i \times q_i \dots\dots\dots [\text{식 5-16}]$$

$$\begin{aligned} & \min_{r_i} P \dots\dots\dots [\text{식 5-17}] \\ & \text{subject to } C = B, \Omega(C, B) \leq \epsilon \end{aligned}$$

소비와 예산이 동일해야 하고, 오차가 허용범위(epsilon) 이하를 만족하면서, 단위당 단가를 최소로 만들어야 한다. 퍼포먼스 광고의 경우 다음과 같은 최적화 목적 함수가 도출된다.

$$\begin{aligned} & \min_{r_i} \Omega(C, B) \dots\dots\dots [\text{식 5-18}] \\ & \text{subject to } P \leq G, B - C \leq \epsilon \end{aligned}$$

퍼포먼스 목표보다 단위당 단가가 낮아야 하며, 예산과 소비의 차이가 허용범위 이하를 만족하면서, 오차를 가장 줄여야만 한다.

특정 성과 목표가 없는 브랜딩 캠페인의 경우는 주요 목표가 예산 소진이다. 그렇기 때문에 각 시간 슬롯 끝에서, 다음 시간 슬롯에 사용될 예산이 정확해지도록 페이싱률을 조정한다.

다음 시간 슬롯에서 사용될 예산은 현재 예산 사용 상황에 따라 결정된다. 주어진 광고 캠페인의 총 예산이 B라고 하면, 예산 사용 계획은 $B = (B(1), \dots, B(k))$ 이고, m 개의 시간 슬롯이 남아있다면, 남은 예측 값들의 합이 남은 예산이 되도록 만든다. 결국 최종 목적 함수는 다음과 같다.

$$\begin{aligned} & \operatorname{argmin}_{\hat{C}^{(m+1)}, \dots, \hat{C}^{(K)}} \Omega \dots\dots\dots [\text{식 5-19}] \\ & \text{subject to } \sum_{t=m+1}^K \hat{C}^{(t)} = B_m \end{aligned}$$

남은 예산 제약을 만족하면서, 오차를 최소화하는 남은 예산 분배를 결정하는 것이다.

5.2.3 현대 포트폴리오 이론(Modern Portfolio Theory)

Markowitz, H. (1952)는 포트폴리오 선택에 대한 논문을 발표했다. 현대 포트폴리오 이론은 다음과 같은 가정을 한다.

1. 투자자는 이성적이고 가능하면 언제나 위험을 피하려 한다.
2. 투자자는 그들이 투자한 최대 결과물을 얻으려 한다.
3. 모든 투자자는 기대하는 수익을 극대화시키기 위한 목적을 공유한다.
4. 시장에서 수수료와 세금은 고려하지 않는다.
5. 모든 투자자는 투자 결정을 위한 모든 필요한 정보에 대해 같은 자료와 레벨에 접근할 수 있다.
6. 투자자는 위험도에 무관하게 자금을 조달할 수 있다.

현대 포트폴리오 이론은 리스크를 싫어하는 투자자가 주어진 수준의 위험(분산)에서 최대의 수익을 내는 포트폴리오를 어떻게 구성하는지에 대한 내용이다. 현대 포트폴리오 이론은 여러 투자에 대한 리스크와 수익을 독립적으로 분석하지 않고도, 포트폴리오를 어떻게 구성하느냐에 따라 포트폴리오의 성과에 영향 미친다는 것이다. 이 이론에서 투자자를 위해 주어진 수준의 리스크 대비 최대 기대 수익을 얻기 위해 특정 지표를 이용하여 포트폴리오를 구성할 수 있는 가능성이 있다.

5.3 연구 설계

이전까지 연구되어 왔던 예산 분배 방식의 기본적인 틀은 입찰 확률을 조정하는 확률적 스로틀 기반의 방식이었다. 그러나 이런 방식들은 광고주의 예산만을 고려하지, DSP 입장에서 이윤 극대화는 적용되지 않았다.

본 연구의 방법은 Lee, et al.(2013)에서 제안된 방법과 유사하게 선형최적화 문제를 설정한다. 모든 시간 슬롯에서의 최적화가 아닌 특정 시점에서의 예산 포트폴리오를 최적화 하는 데에 목적을 둔다. 또한 퍼포먼스의 시간대 비율 변화로 캠페인의 성과 목표를 달성하려는 것과 달리 본 연구의 기대 수익은 퍼포먼스 목표를 초과하는 수익으로 정의함으로써, 명시적으로 모형 내에 퍼포먼스 목표를 도입한다. 연구 방법에서 수익률은 유저 반응률의 함수로 정의할 수 있다. 기대 위험은 마찬가지로 기대 수익의 변동성으로 추정한다. Sharpe Ratio를 활용해 리스크 대비 투자 결과를 측정할 것이며, Sharpe Ratio에서 무위험 기대수익은 퍼포먼스 목표로 정한다. 이로써 해당 내용을 반영하면, 캠페인 별 퍼포먼스 목표를 달성하기 위한 제약 사항을 모형식에 포함하게 된다. 퍼포먼스 목표가 없는 경우에는 0으로 처리하게 된다. 학습 자료를 이용해 각 시간대별 퍼포먼스를 측정한 후 Sharpe Ratio가 최대가 되는 가중치를 사용해서 입찰 확률로 사용할 것이다.

$$Sharpe Ratio = \frac{R_p - R_{rf}}{\sigma_p} \dots\dots\dots [식 5-20]$$

R_p = 포트폴리오의 기대이익
 R_{rf} = 무위험수익
 σ_p = 포트폴리오의 표준편차

균등 분배 방식과 페이싱룰 방식 그리고 포트폴리오 방식을 비교할 것이다.

5.4 결론

5.4.1 학습 결과

입찰 전략은 일반적으로 상수입찰, 랜덤입찰, 최대 CPC(Cost Per Click) 입찰, 선형 입찰이 있다. 하지만 Zhang, et al.(2014)에서는 선형 입찰방식이 벤치마킹에서 최적임을 보고하고 있다. 따라서 본 연구는 예산분배에 관한 연구이기 때문에 입찰 전략은 고정된 것으로 보고 선형 입찰 전략을 사용해 각각 예산 분배 방식을 비교한다.

선형 입찰 방식은 입찰 가격을 예측 CTR(Click Through Rate)의 선형적 비율로 하는 것이다. 공식은 일반적으로 다음과 같다.

$$\text{입찰가격} = \text{기본입찰가} \times \text{예측}CTR / \text{평균}CTR \dots\dots\dots [\text{식 } 5-21]$$

평균 CTR은 학습 자료를 통해 얻어지고, 기본입찰가는 튜닝이 가능한 값이다. 본 연구에서는 1부터 800까지 기본 입찰가를 실험하였다. 평가 자료에서 하루 총 광고비를 예산으로 주어진 것으로 보았다.

[표 5-1] 예산 분배의 결과 (클릭 수)

광고주코드	시즌	균등분배	페이싱룰	포트폴리오
1458	2	497	326	316
2259	3	83	123	121
2261	3	51	94	94
2821	3	262	376	368
2997	3	349	530	516
3358	2	205	90	87
3386	2	351	362	352
3427	2	343	238	231
3476	2	273	240	232

시즌 2에서는 균등분배가 더 낮고, 시즌 3의 자료에서는페이싱률이나 포트폴리오의 클릭 수가 더 많다.

[표 5-2] 예산 분배의 결과 (예산소진율)

광고주코드	시즌	균등분배	페이싱률	포트폴리오
1458	2	59.29%	58.24%	54.56%
2259	3	64.47%	93.25%	89.33%
2261	3	40.80%	96.13%	93.34%
2821	3	63.82%	99.38%	94.71%
2997	3	58.06%	99.70%	95.17%
3358	2	10.83%	33.32%	32.44%
3386	2	51.12%	79.37%	75.51%
3427	2	39.49%	63.29%	60.24%
3476	2	33.11%	75.10%	71.25%

[표 5-3] 예산 분배의 결과 (기본 입찰가)

광고주코드	시즌	균등분배	페이싱률	포트폴리오
1458	2	727	800	778
2259	3	284	800	798
2261	3	159	800	795
2821	3	182	800	800
2997	3	188	800	699
3358	2	114	800	799
3386	2	275	800	793
3427	2	297	800	790
3476	2	193	800	798

[표 5-4] 예산 분배의 결과 (CTR)

광고주코드	시즌	균등분배	페이싱률	포트폴리오
1458	2	0.109%	0.071%	0.073%
2259	3	0.025%	0.030%	0.031%
2261	3	0.021%	0.028%	0.029%
2821	3	0.049%	0.057%	0.058%
2997	3	0.317%	0.340%	0.346%
3358	2	0.390%	0.070%	0.070%
3386	2	0.102%	0.075%	0.076%
3427	2	0.129%	0.060%	0.061%
3476	2	0.117%	0.054%	0.055%

[표 5-5] 예산 분배의 요약표

모델	예산소진을 향상률(%)	CTR 향상률(%)
균등분배		
페이싱률	65.75%	-37.65%
포트폴리오	58.33%	-36.54%

균등분배는 확실히 전체 예산을 다 소진 못하는 것이 보인다. 광고주는 목적은 2가지이다. 부드럽게 예산을 소진하면서 많은 사람에게 노출시키는 것과 성과 지표를 맞추는 것이다. 그런데 CTR은 노출과 상충관계에 있다. 분모가 노출이기 때문에 노출을 줄이면, CTR은 오르게 되어있다. 그래서 예산 분배는 예산 소진과 퍼포먼스 사이의 균형을 잘 맞춰야한다. 균등분배는 예산 소진을 거의 못하기 때문에 베이스라인으로 쓰기엔 좋지만, 부드러운 게재 입장에서는 좋지 않다. [표 5-5]에서 균등분배의 CTR이 매우 좋게 나오는데 이는 제한된 예산 때문에 노출을 제대로 못했기 때문이다. 3358 광고주의 경우 대부분의 클릭이 0시에 몰려있기 때문에 균등분배의 CTR이 매우 높게 나오게 된다. 결과적으로 예산 소진을 대부분 하면서도, DSP 입장에서의 기대수익인 CTR을 높게 맞추는 포트폴리오 방식이 페이싱률 방식을 약간 개선했다고 볼 수 있다.

본 연구는 예산 분배에 관한 연구이다. 이전까지의 온라인 디스플레이

광고에서 예산 분배는 입찰 가격을 조정하여 낙찰 확률을 간접적으로 조정하는 방식인 입찰가 조정 방식과 입찰 확률을 조정하여 예산을 분배하는 확률적 스로틀 방법으로 연구되고 있었다. 그러나 확률적 스로틀 방법의 경우 해 공간의 크기가 크고, 시장의 동태적 변화와 다중 목적 함수 및 다중 제약을 가진 최적화 문제로서 풀기 어려웠다.

온라인 디스플레이 광고 시장은 각 DSP를 중개회사로 보고, 거래되는 광고 게재 기회는 주식으로 보면, 주식 시장과 매우 닮아있다.

본 연구에서는 현대 포트폴리오 이론을 활용하여, 입찰 확률에 대한 시뮬레이션을 통해 기대 수익을 극대화하는 방법을 논의했다. 최적 포트폴리오로 구성된 입찰 확률을 사용하여 예산 분배를 할 수 있다.

Ⅵ. 결 론

광고란 서비스를 구독하거나 상품을 구매하려는 잠재적인 고객을 대상으로 하는 마케팅 기법이다. 광고는 또한 브랜드와 관련된 광고의 반복된 노출을 통해 브랜드 이미지를 만드는 방법이기도 하다. 기술의 발전과 함께 사람들의 생활 반경은 오프라인에서 온라인으로 옮겨졌다. 온라인 광고 산업의 핵심 목표는 사용자 경험을 바탕으로 사용자의 인터넷 경험을 향상시키고, 광고주의 광고 효과를 높이며, 매체의 이윤을 극대화하는 것이다. 이를 위해 쿠키를 이용하여, 브라우저의 행동 정보를 수집한다. 그리고 수집된 정보를 바탕으로 브라우저의 사용자 프로파일링을 한다. 그런 다음 프로파일링 한 정보를 바탕으로 개인화된 광고를 내보낸다. 이런 개인화 광고를 내보내기 위해서 시스템이 필요해졌다.

온라인 광고 시장은 처음 단순한 광고주와 매체의 연결만 있었지만, 시간이 지남에 따라 이를 대행해주는 회사들이 생겼다. 또한 그런 대행해주는 회사를 중개하는 네트워크 회사들이 점차 커지면서, 거대한 시장이 되었다. 이런 시장에서 상품은 특정 영역에 광고를 송출할 수 있는 기회가 되기 때문에 시간 소멸성 상품이다. 이런 시장에서 거래 효율성을 높이기 위해서는 경매 방식이 효율적이게 된다.

온라인 광고 시장은 두 번째 가격 경매 방식을 이용한다. 이 방식은 광고주로 하여금 자신의 사적 가치를 스스로 드러내게 하는 방식이기 때문에 효율적이다. 사적 가치는 클릭 확률과 동일하다고 할 때, 클릭 확률을 정확하게 추정하는 것은 기대 이익과 직접적으로 연관이 있다. 본 연구에서는 클릭 확률 추정 모델을 만들 때에 베이스 모델로는 FM-FTRL을 사용했다. 또한 추가적인 변수로 입력 벡터의 비선형 그룹을 찾아내는 K-평균 클러스터링 알고리즘의 예측치를 사용했다. 그리고 그라디언트 부스팅 모델의 잎 노드들을 추가적인 변수로 사용했다. K-MEANS 변수와 GBDT 변수가 각각 원 자료의 비선형 특성을 전체 자료의 10%만 가지고도 잘 포착하며, 일반화된다고 볼 수 있다. 광고주 코드별로 LogLoss나 RMSE가 차이나는 것은 산업 특성으로 보인다. 시즌 2에서 K-MEANS 변수와 GBDT 변수가

성능이 안 좋은 이유는 K-MEANS 변수와 GBDT 모두 입력 변수들의 비선형 그룹을 찾는 방식인데, 그 방식이 iPinYou가 분류해놓은 유저 태그와 충돌(Collision)이 일어나는 것으로 추정된다. iPinYou가 자신들의 DAAT를 이용해 유저 태그를 만들 때, 데모그래픽, 지오그래픽, 구매 성향 등의 여러 자료를 사용하는 것이 K-MEANS 같은 방식과 유사한 방식으로 그룹핑 된 것이라면, 본 학습 모델인 FM-FTRL에서 변수 독립이 되지 않아 계수 값을 나눠가짐으로써 성능이 저하된 것일 수도 있다.

예산 분배는 온라인 광고 시장에서 또 하나의 중요한 문제이다. 광고주는 시간에 따라서 부드럽게 게재되는 예산 할당을 추구한다. 시간에 따라서 부드럽게 예산 분배가 된다면, 같은 광고비를 지속적으로 사용할 수 있고, 더 넓은 잠재고객에게 도달할 수 있기 때문이다. DSP는 개별 증권사와 매우 유사하다. 증권사에서 거래하는 주식은 광고 게재 기회가 되고, 개별 주체가 서로 다른 주식을 거래하며, 기대 이윤을 극대화하기 위해 많은 작업을 한다. DSP를 마치 증권사처럼 생각한다면, 예산 분배는 포트폴리오를 작성하는 것이라 생각할 수 있다. 광범위한 주식 중에서 최적 기대 이윤을 가져다 줄 최적 포트폴리오를 결정하는 문제와 동일시 될 수 있다. 일반적인 예산 분배 방법은 스로틀링과 입찰 수정 방법이 있다. 본 연구에서는 스로틀링 방식을 사용했다. 개선점으로는 DSP 입장에서 포트폴리오 이론을 바탕으로 한 최적 입찰 확률을 추정하였다. Sharpe Ratio를 활용해 리스크 대비 투자 결과를 측정하고 최대 비율이 되는 가중치를 사용하여, 예산 소진을 거의 다 하면서도 퍼포먼스를 향상시키는 방안을 제시하였다.

본 논문은 온라인 광고 시장에서 경제학을 활용하는 것을 주제로 하여, 핵심문제들에 대한 종합적인 분석을 진행한 연구이다. 온라인 광고 시장은 그 시장 규모나 발달 속도에 비해 우리나라에 알려진 바가 적고, 해외에서는 활발히 연구되고 있다. 본 연구는 온라인 광고 시장에 핵심 문제들에 대한 최신 알고리즘에 대한 개괄적인 소개와 더불어 개선할 수 있는 방향을 제시하는 방안을 제안함으로써 계량경제학 분야에 학술적인 기여를 할 수 있을 것으로 기대한다. 온라인 광고 시장은 빅 데이터라 불리는 방대한 양의 데이터를 가지기 때문에 스몰 데이터에서는 문제가 되지 않던 것들이 문제가

되는 경우가 많다. 단순히 양이 많아짐에 따라 발생하는 문제들을 처리하는 방법 또한 잘 알려져 있지 않고, 구전되는 경우가 많기 때문에 해당 내용에 대한 논의를 통해 실무에 도움이 되기를 기대한다.

참 고 문 헌

1. 국내문헌

- 이현채. (2012년 1월 11일). [Ads. Study] Display Ads. 3편 - 초기
광고의 거래(?) 형태와 진화, 검색일 : 2018년 10월 20일,
검색 주<http://itmobile.tistory.com/entry/AdsStudyDisplayAds>
- 이현채. (2012년 1월 18일). [Ads. Study] Display Ads. 4편 - AD
Network의 역할과 BM, 검색일 : 2018년 10월 20일,
<http://itmobile.tistory.com/entry/ADNetwork>
- 이현채. (2012년 1월 27일). [Ads. Study] Display Ads. 5편 - AD
Exchange의 역할, 검색일 : 2018년 10월 20일,
<http://itmobile.tistory.com/entry/ADExchange>

2. 국외문헌

- Borgs, C., Chayes, J., Immorlica, N., Jain, K., Etesami, O., & Mahdian, M. (2007, May). Dynamics of bid optimization in online advertisement auctions. In *Proceedings of the 16th international conference on World Wide Web* (pp. 531-540). ACM.
- Chen, Y., Berkhin, P., Anderson, B., & Devanur, N. R. (2011, August). Real-time bidding algorithms for performance-based display ad allocation. In *Proceedings of the 17th ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 1307-1315). ACM.
- Cox, D. R. (1958). The regression analysis of binary sequences. *Journal of the Royal Statistical Society. Series B (Methodological)*, 215-242.

- Edelman, B., & Ostrovsky, M. (2007). Strategic bidder behavior in sponsored search auctions. *Decision support systems*, 43(1), 192–198.
- Friedman, J. H. (2001). Greedy function approximation: a gradient boosting machine. *Annals of statistics*, 1189–1232.
- He, X., Pan, J., Jin, O., Xu, T., Liu, B., Xu, T., ... & Candela, J. Q. (2014, August). Practical lessons from predicting clicks on ads at facebook. In *Proceedings of the Eighth International Workshop on Data Mining for Online Advertising* (pp. 1–9). ACM.
- Karlsson, N., & Zhang, J. (2013, June). Applications of feedback control in online advertising. In *American Control Conference (ACC), 2013* (pp. 6008–6013). IEEE.
- Lee, K. C., Jalali, A., & Dasdan, A. (2013, August). Real time bid optimization with smooth budget delivery in online advertising. In *Proceedings of the Seventh International Workshop on Data Mining for Online Advertising* (p. 1). ACM.
- Lloyd, S. (1982). Least squares quantization in PCM. *IEEE transactions on information theory*, 28(2), 129–137.
- Markowitz, H. (1952). Portfolio selection. *The journal of finance*, 7(1), 77–91.
- McMahan, H. B., Holt, G., Sculley, D., Young, M., Ebner, D., Grady, J., ... & Chikkerur, S. (2013, August). Ad click prediction: a view from the trenches. In *Proceedings of the 19th ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 1222–1230). ACM.
- Mehta, A., Saberi, A., Vazirani, U., & Vazirani, V. (2007). Adwords and

- generalized online matching. *Journal of the ACM (JACM)*, 54(5), 22.
- Menezes, F. M., & Monteiro, P. K. (2005). *An introduction to auction theory*. OUP Oxford.
- Rendle, S. (2010, December). Factorization machines. In Data Mining (ICDM), *2010 IEEE 10th International Conference on* (pp. 995–1000). IEEE.
- Ta, A. P. (2015, October). Factorization machines with follow-the-regularized-leader for CTR prediction in display advertising. In *Big Data (Big Data), 2015 IEEE International Conference on* (pp. 2889–2891). IEEE.
- Tibshirani, R. (1996). Regression shrinkage and selection via the lasso. *Journal of the Royal Statistical Society. Series B (Methodological)*, 267–288.
- Wang, J., Zhang, W., & Yuan, S. (2016). Display advertising with real-time bidding (RTB) and behavioural targeting. *arXiv preprint arXiv:1610.03013*.
- Weinberger, K., Dasgupta, A., Attenberg, J., Langford, J., & Smola, A. (2009). Feature hashing for large scale multitask learning. *arXiv preprint arXiv:0902.2206*.
- Wolpert, D. H. (1992). Stacked generalization. *Neural networks*, 5(2), 241–259.
- Xu, J., Lee, K. C., Li, W., Qi, H., & Lu, Q. (2015, August). Smart pacing for effective online ad campaign optimization. In *Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp.

2217–2226). ACM.

- Zhang, W., Rong, Y., Wang, J., Zhu, T., & Wang, X. (2016, February). Feedback control of real-time display advertising. In *Proceedings of the Ninth ACM International Conference on Web Search and Data Mining* (pp. 407–416). ACM.
- Zhang, W., Yuan, S., Wang, J., & Shen, X. (2014). Real-time bidding benchmarking with ipinyou dataset. *arXiv preprint arXiv:1407.7073*.
- Zhang, W., Zhou, T., Wang, J., & Xu, J. (2016, August). Bid-aware gradient descent for unbiased learning with censored data in display advertising. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining* (pp. 665–674). ACM.

ABSTRACT

Estimation of Click Probability and Budget Allocation in the Online Advertising Market

Hae, Hyeon-yong

Major in Economics & Real Estate

Dept. of Economics

The Graduate School

Hansung University

Advertising is a marketing technique that targets potential customers who want to subscribe to a service or purchase a product. Advertising is also a way to create a brand image through repeated impressions of ads related to the brand. With the advancement of technology, people's lives have been moved from offline to online. A key goal of the online advertising industry is to improve the user experience on the Internet based on the user experience, to increase the advertising effectiveness of the advertiser, and to maximize the media profit. To do this, we use cookies to collect browser behavior information. And then profiling the user of the browser based on the collected information. It then sends personalized ads based on the information you profiled. We needed a system to export these personalized ads.

This paper is a study on algorithms in this online advertising

system. First of all, I will give an overview of the online advertising market. After that, I discuss the economic characteristics of the online advertising market and present major problems. The online advertising market is less known in Korea than the market size and development rate, and has been actively studied overseas. This study is expected to contribute to academic economics in the field of econometrics by proposing an overview of the latest algorithms on the core problems in the online advertising market and suggesting ways to improve them.

First, I introduce the algorithms related to the click probability estimation which is the main problem of the online advertisement market and discuss the improvement plan. In this study, FM-FTRL was used as the base model when creating the click probability estimation model. I also used the predictions of a K-means clustering algorithm to find nonlinear groups of input vectors as additional variables. And I used the leaf nodes of the gradient boosting model as additional variables. The K-MEANS variable and the GBDT variable capture the nonlinear characteristics of the original data well, even with only 10% of the total data, and are generalized.

Next, I discuss budget allocation, another important issue in the online advertising market. Advertisers want a budget allocation that will appear smooth over time. Given the smooth distribution of budget over time, it can continue to use the same advertising costs and reach a wider audience. DSPs are very similar to individual securities firms. If we think of DSP as a brokerage, we can think of the budget allocation as creating a portfolio. It can be identified with the problem of determining an optimal portfolio that will yield optimal expected profit among a broad range of stocks. Common budget allocation methods include throttling and bidding modifications. The throttle method was used in this study. As an improvement, I estimate the optimal bid probability based on the

portfolio theory in the DSP perspective. The Sharpe Ratio was used to measure the return on investment and to use the maximum weighting value. And suggested ways to improve performance while exhausting the budget.

Because the online advertising market has a vast amount of data called big data, things that are not problematic for small data are often problematic. It is not well known how to deal with problems that arise as the quantity simply increases, so I expect to be helpful to the practitioner by discussing the contents.

【Keyword】 Large-scale Data, Online Learning, Online Advertising, Stacked Generalization, Budget Allocation, CTR Prediction