

투고일 : 2019.12.30

심사일 : 2020.1.22

게재확정일 : 2020.1.29

2

치의학 연구에서 R program을 이용한 성향점수매칭의 단계적 안내

¹⁾전남대학교 치과병원²⁾전남대학교 치의학전문대학원 치과교정학교실, 치의학 연구소안화연¹⁾, 임희정²⁾

ABSTRACT

A step-by-step guide to Propensity Score Matching method using R program in dental research

¹⁾Chonnam National University Hospital²⁾Department of Orthodontics, Chonnam National University, School of Dentistry, Dental Science Research InstituteHwayoen An¹⁾, Hoi-Jeong Lim²⁾

The propensity score matching method is a statistical method used to reduce selection bias in observational studies and to show effects similar to random allocation. There are many observational studies in dentistry research, and differences in baseline covariates between the control and case groups affect the outcome. In order to reduce the bias due to confounding variables, the propensity scores are used by equating groups based on the baseline covariates. This method is effective, especially when there are many covariates or the sample size is small.

In this paper, the propensity score matching method was explained in a simple way with a dental example by using R software. This simulated data were obtained from one of retrospective study. The control group and the case group were matched according to the propensity score and compared before and after treatment. The propensity score matching method could be an alternative to compensate for the disadvantage of the observation study by reducing the bias based on the covariates with the propensity score.

Keywords: Propensity Score Matching, matched case-control study, dental research

Corresponding Author

Hoi-Jeong Lim, PhD

Department of Orthodontics, Chonnam National University, School of Dentistry, Dental Science Research Institute
33 Yongbong-ro, Buk-gu, Gwangju 500-757, South Korea

Tel: +82-62-530-5830 Fax: +82-62-530-5659 E-mail: hylim@jnu.ac.kr

ACKNOWLEDGMENT This Research was supported by the National Research Foundation of Korea (NRF) grant, funded by the Korea government (No. 2018R1D1A1B07049719).

I. 서론

1920년대 말 홍차의 나라 영국에서 밀크티에 정통한 한 부인이 자신은 '밀크티에 홍차를 먼저 넣었는지, 우유를 먼저 넣었는지' 구분할 수 있다고 주장했다. 그 당시 그들이 배운 과학적 지식에 근거하면 홍차와 우유를 섞었을 때 어떤 것을 먼저 넣든지 화학적 성질의 차이가 없다고 알려져 있었다. 그때 현대 통계학의 아버지인 피셔가 홍차를 먼저 넣은 밀크티 네 잔, 우유를 먼저 넣은 밀크티 네 잔, 총 여덟 잔을 준비한 후 모든 잔을 무작위로 섞었다. 그리고 그 부인으로 하여금 우유를 먼저 넣은 밀크티 네 잔을 모두 고르도록 하여 정말 밀크티에 무엇을 먼저 넣었는가를 구분할 수 있는지 알아보는 실험을 하였다. 그녀는 1/70의 우연이라고 하기는 힘든 확률을 뚫고 모두 맞췄다고 한다. 2003년에 와서야 영국 왕립 화학협회에서 뜨거운 홍차에 우유를 넣으면 우유 단백질이 변성되기 때문에 차이가 있다는 보도를 발표하였다. 피셔의 실험은 간단한 것 같지만 이는 '과학적으로 실증하기 위한 순서' 중에서 '임의(무작위)'라는 개념을 적용한 것으로 현재 연구 과정 중에서 가장 중요한 부분이라고 할 수 있다.¹⁾

현재까지 여러 가지 통계 방법들이 소개되었고²⁻⁵⁾, 여기서 소개하고자 하는 성향점수 매칭방법(propensity score matching method: PSM)을 통해 임의라는 개념이 왜 중요하게 여겨지고 있는지 알 수 있다. 어떠한 처치에 대한 효과를 알아보고자 하는 고전적인 연구방법으로 전향적인 연구(prospective study) 방법과 관측연구(observational study) 방법이 있다. 전향적인 연구방법은 특정 처치가 나타내는 결과를 추정하기 위해 연구대상을 무작위 배정하여 결과에 영향을 주는 특성의 차이가 없도록 하는 방법이다. 그러나 무작위배정이 임상에서 윤리적인 문제를 일으킬 수 있다. 반면에 관측연구방법은 무작위배정 없이 특정 집단을 대상으로 연구를

진행하기 때문에 임상에 적용이 쉽고 특정 교란변수에 대한 영향을 통계적으로 배제 시킬 수 있다. 그러나 연구대상 선정 시 선택편향을 피할 수 없어 어떤 현상의 인과관계를 추론하는 것은 불가능하다.

이러한 고전적인 연구 방법들이 가지는 단점을 극복하고 선택편향을 최소화하기 위해 짝짓기(matching) 방법을 사용하는데 일반적인 짝짓기방법은 발생률이 낮은 연구대상의 경우 통계적 검정력을 가질 만큼 표본을 수집하는 것이 어렵고, 수집된다 하더라도 왜곡된 결과를 도출할 수 있다. 그렇기 때문에 일반적인 짝짓기 방법보다는 보다 정확하고 손쉬운 성향점수 매칭방법이 사용된다. PSM은 실험군과 대조군의 공변량의 영향력을 반영하는 성향점수를 계산하여 피험자를 짝짓기 함으로써 랜덤하게 배정된 효과를 주어 대조군을 재구성 함으로써 하고자 하는 연구결과의 신뢰도와 정확도를 향상시킬 수 있기 때문에 관측연구 뿐 아니라 무작위 배정이 어려운 연구 등에서 선택편향을 줄이기 위해 사용되고 있다. PSM은 의학계에서 무작위 대조 시험(Randomized Controlled Trial)이 어려운 상황에서 치료효과를 알아보고자 고안되었지만 선택편향을 줄이고 임의화를 실현시킬 수 있을 정도로 기저특성차이를 더 잘 조정할 수 있고, 다른 가정이 필요하지 않기 때문에 요즘 연구 분야에서 많이 사용되고 있다.

본 논문에서 PSM방법의 소개와 치의학 분야에서의 R program을 이용한 PSM방법의 적용과정을 예제를 이용하여 보여줌으로써 치의학에서 더 실용적으로 연구에 이용되어 올바른 결과 도출에 도움이 되고자 한다.

II. 성향점수매칭에 사용된 예제

후향적 연구로부터 얻어진 데이터를 가상으로 만들어 성향점수매칭을 설명하였다. 실험군은 치과 교정 치

료의 일부로 급속상악확장장치(Rapid maxillary expansion) 치료를 받은 환자로 가정하였고 대조군은 급속상악확장장치 없이 일반 치과 교정 치료를 받은 환자들로 가정하였다. 급속상악확장장치 치료를 받기 전 처음 CBCT scan을 촬영한 시기를 T0로 하고 최소 3개월이 지난 후에 다시 CBCT scan을 촬영한 시기를 T1이라고 하였다. 그때의 airway MCA(minimum cross-sectional area)를 측정하여 급속상악확장장치가 기도 부피(airway volume)를 증가시키는지 알아보려고 3D CT를 이용하여 작성된 가상의 데이터이다. 총 50명의 데이터 중에 실험군 17명과 대조군 33명을 각각 대상으로 하였다.

III. 성향점수매칭 과정

A. R program을 설치하고 실행하는 과정

이 논문의 분석은 무료 프로그램인 R studio를 기반으로 했지만 R studio를 실행시키기 위해서는 R 프로그램이 먼저 설치되어야 한다. 아래의 과정은 R 3.3.3을 기반으로 실행하였다. <https://cran.r-project.org/>에 들어가서 R을 설치한 후 <https://www.rstudio.com/products/rstudio/download/>에 들어가서 R studio를 설치한다. 설치가 끝나면 R studio에 들어가서 메뉴바에서 Tools>Install Packages를 선택한다. Install Packages의 윈도우가 뜨면 Packages 항목에 MatchIt, optmatch를 각각 적고 install을 누르면 두 개의 패키지가 설치된다. 설치 후 두 함수 항목을 체크하면 이 패키지가 로딩된다 (Fig. 1).

B. 데이터를 불러오는 과정

사용된 변수는 group(0:RME를 사용하지 않은 그룹(대조군), 1:RME를 사용한 그룹(실험군))이고 공변량 변수로 사용된 변수는 4개로, gender(0:남, 1:여), age, CBCT(다음 CBCT촬영까지의 개월 수), Class(1,2,3: 교합관계)로 정의되었다. 결과 변수로 MCA-T0(급속상악확장장치 치료 전의 airway volume), MCA-T1(급속상악확장장치 치료 후의 airway volume)으로 정의되었다. 아래의 데이터는 전남대학교 치의학전문대학원(<http://dent.jnu.ac.kr>) 일반대학원 자료실에서 번호 30번 Propensity score matching data를 클릭하여 matchdata를 다운받을 수 있다.

분석할 data를 엑셀에서 정리한 후 확장명을 반드시 csv로 저장한다(Fig 2). 저장된 데이터를 R studio로 불러오기 위하여 아래의 명령어를 입력하면 파일을 선택할 수 있는 창이 뜨고 저장된 엑셀파일 50 patients를 선택한다(Fig 3). "data1"이라는 이름으로 R studio에 저장될 것이다.

```
>data1 = read.csv(file.choose())
```

제대로 불러왔는지를 확인하기 위해서 "data1"을 입력하면 데이터가 보인다(Fig. 4).

```
>data1
```

변수명만으로 데이터에 대한 정보를 불러올 수 있도록 R studio에 데이터를 attach하는 과정이 필요하다.

```
>attach(data1)
```

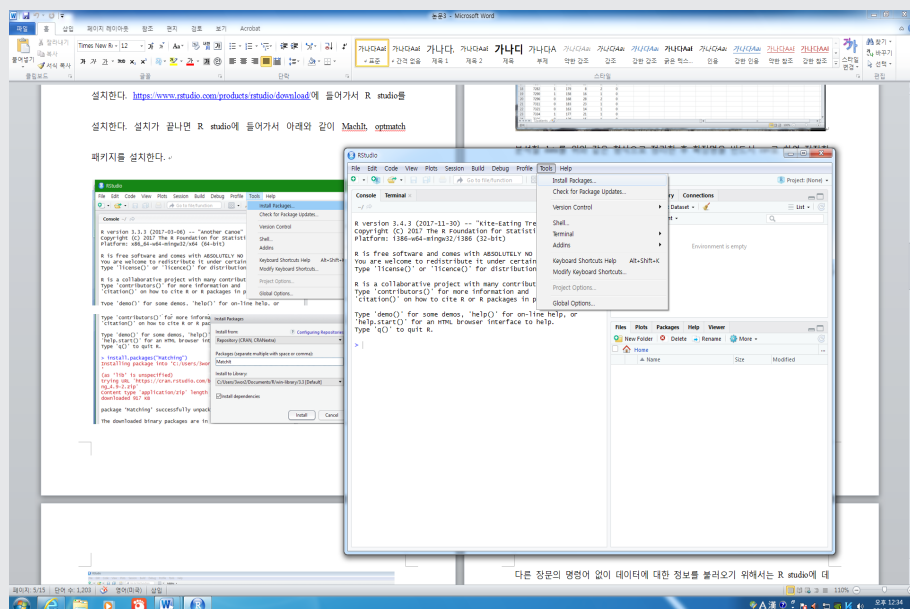


Fig. 1. R program 패키지 설치과정

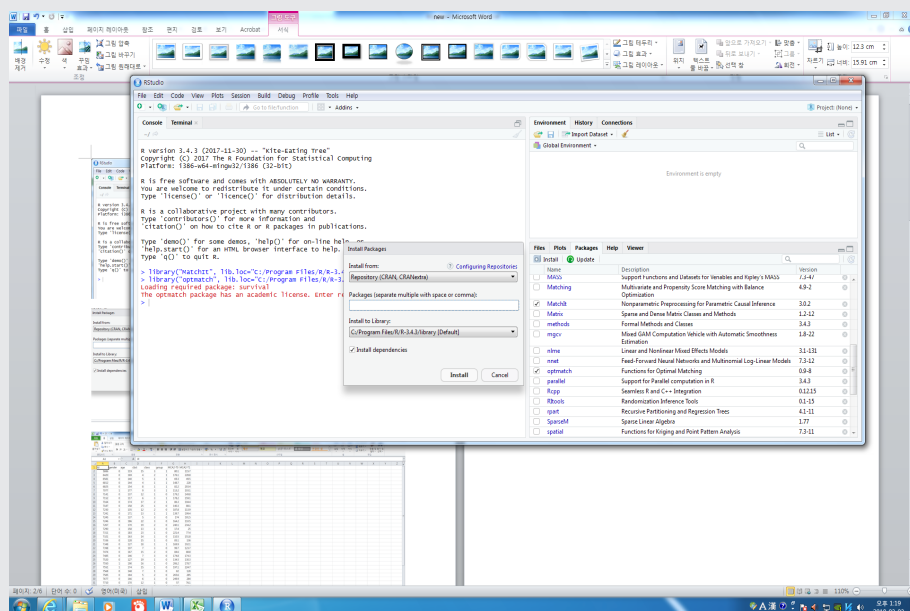


Fig. 2. 사용된 데이터

A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U	V	W	X	Y	Z
1	gender	age	cbc	class	group	MCA-70	MCA-71																		
2	7141	0	137	12	1	0	179.2	149.8																	
3	7290	1	158	13	1	0	17.4	25																	
4	7311	0	183	23	1	0	221.4	77.4																	
5	7321	0	163	14	1	0	110.5	151.8																	
6	7398	0	107	7	3	0	99.7	123.7																	
7	7474	0	167	15	2	0	84.6	88.8																	
8	7485	0	146	7	3	0	174.8	174.3																	
9	7520	0	127	19	1	0	134.5	130.3																	
10	7551	1	174	15	1	0	191.1	224.7																	
11	7585	0	184	5	2	0	265.6	285																	
12	7677	0	166	6	1	0	249.9	284																	
13	7710	0	170	12	1	0	57	76.1																	
14	7711	0	190	10	1	0	126.2	183.4																	
15	7767	0	170	9	2	0	169.1	158.6																	
16	7789	1	153	6	3	0	82.7	108																	
17	7814	0	163	7	1	0	254.2	267.2																	
18	7844	1	185	17	3	0	182.5	236.5																	
19	7921	0	187	7	2	0	141.9	126.3																	
20	7945	1	132	10	3	0	43.5	19.4																	
21	7971	0	144	5	1	0	92.3	104.1																	
22	8002	0	169	9	3	0	79.9	83																	
23	8099	0	186	18	2	0	182.9	223.1																	
24	8045	0	163	7	1	0	259.3	303.7																	
25	8139	1	154	14	2	0	49.9	75.9																	
26	8171	0	147	16	2	0	78.2	97.8																	
27	8225	0	151	5	2	0	111.8	53.6																	
28	8277	0	160	12	2	0	25.7	31.2																	
29	8283	0	131	13	1	0	90.9	112.7																	
30	8326	0	116	5	3	0	67.2	73.5																	
31	8435	0	140	7	3	0	334.2	322.2																	
32	8440	0	159	15	1	0	143.2	148.6																	
33	8486	0	142	5	3	0	269	308.4																	
34	8529	1	140	15	1	0	67.7	53.2																	
35	8586	0	119	15	3	1	80.1	115.7																	
36	8420	0	169	4	2	1	176.1	228.8																	
37	6981	0	140	5	1	1	69.3	49.5																	

Fig. 3. R studio에서 파일을 불러오는 과정

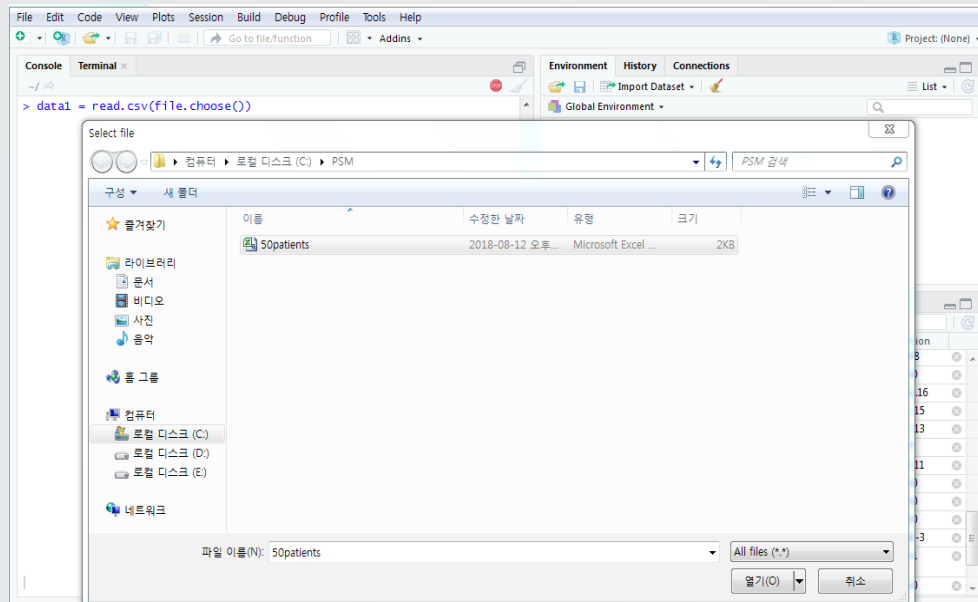


Fig. 4. R program에 보이는 데이터

C. 성향점수매칭의 실행

R program을 수행하기 전에 성향점수를 계산하는 방법과 매칭방법의 종류는 다음과 같다. 성향점수는 공변량이 주어졌을 때, 실험군 또는 대조군에 들어갈 확률 예측값으로 로지스틱 모형을 이용하여 추정할 수 있다.

이렇게 추정된 확률이 성향점수이고 실험군과 대조군에서 성향점수가 같은 피험자끼리 매칭하면 랜덤하게 배정된 것과 같은 효과를 가질 수 있다.

매칭방법에는 최근접이웃짜짓기 (Nearest neighbor matching), 반경짜짓기 (Radius matching), 범위짜짓기 (Caliper matching), 최적짜짓기 (Optimal matching) 방법으로 크게 4가지가 있다^{6,7)}. 최근접이웃짜짓기 방법은 대조군과 실험군에 포함된 모든 연구 대상들의 추정된 성향점수차이의 값이 가장 작은 순서로 짜짓기를 하는 방법이다. 반경짜짓기 방법은 실험군의 성향점수로부터 미리 설정한 간격 이내의 범위에 들어오는 대

조군 중에 가장 가까운 개체를 선택하는 방법이다. 대상들 사이의 성향점수 차이가 이 범위 내에 해당되는 경우에만 짝을 이루어 분석에 포함시키고 제외되는 모든 개체는 분석에서 제외하는 방법이다. 이 방법은 대조군의 수가 충분하다면 간격이 좁을수록 정확성이 좋다. 범위짜짓기 방법은 반경짜짓기의 한 종류로 실험군의 추정된 성향점수의 표준오차의 1/4에 해당되는 값을 범위로 지정하는 방법이다. 최적짜짓기 방법은 유사한 성향점수를 가진 대조군과 실험군의 연구대상들이 하나의 계층으로 분류하여, 자료전반에 걸쳐 층화를 시행하는 방법이다. 여러 개의 계층으로 나누어지는 과정에서 짜짓기가 이루어지며, 각 계층 내에서 실험군과 대조군 표본수의 비율에 따라 matching process가 결정된다. 최근접이웃짜짓기 방법은 유사하게 보이나 최적짜짓기 방법과 다르게 국소적으로만 최적의 짜짓기가 시행된다.

R program에서는 최적짜짓기 방법을 사용하는데 최근접이웃짜짓기 방법과의 차이는 간단한 예를 보면 이

$$y_i = \ln\left(\frac{p_i}{1-p_i}\right) = \alpha + \beta_1 X_{1i} + \beta_2 X_{2i} + \cdots + \beta_p X_{pi} \quad \text{-- 로지스틱 모형}$$

$$p_i = \Pr(Y=1|x) = \frac{e^{\alpha + \beta_1 x_{1i} + \beta_2 x_{2i} + \cdots + \beta_p x_{pi}}}{1 + e^{\alpha + \beta_1 x_{1i} + \beta_2 x_{2i} + \cdots + \beta_p x_{pi}}} = \frac{1}{1 + e^{-(\alpha + \beta_1 x_{1i} + \beta_2 x_{2i} + \cdots + \beta_p x_{pi})}}$$

Table 1. optimal 매칭방법과 nearest neighbor 매칭방법의 비교에 사용된 예제

실험군	Propensity score	대조군	Propensity score
A	5.7	V	5.5
B	4.0	W	5.3
C	3.4	X	4.9
D	3.1	Y	4.9
		Z	3.9

해할 수 있을 것이다(Table 1).

위의 두 그룹을 매칭 한다고 하면 최근접이웃짝짓기 방법은 전반적인 데이터를 고려하지 않고 첫 번째 데이터부터 차례대로 가장 비슷한 숫자의 데이터와 매칭된다. 따라서 실험군 A와 대조군 V가 매칭되고 차례대로 B-Z, C-X, D-Y로 짝지어지게 되어 가장 마지막 데이터가 매칭되었을 때 점수 차이가 가장 많이 발생할 수 있다. 반면에 최적짝짓기 방법은 전반적인 데이터를 고려하여 매칭되어 나가기 때문에 A-V, B-X, C-Y, D-Z로 점수 차이가 전반적으로 비슷하게 나타날 수 있다.

다음은 MatchIt package의 “matchit” 명령어를 사용하여 성향점수매칭을 수행하는 과정이다.

```
>matchingdata = matchit(group ~ gender +
  age + cbct + class, data = data1, method =
  'optimal', ratio = 1)
```

Gender, age, cbct 와 class는 매칭 변수들이고 group은 실험군 또는 대조군을 나타내는 변수이다. Method는 여러 매칭 방법들 중에서 하나인 최적짝짓기 방법을 사용하였다. 매칭된 데이터에 대한 정보를 불러 오기 위해서 아래와 같은 명령어를 사용한다.

```
>summary(matchingdata)
```

summary 명령어를 통해 실험군과 대조군에서의 매칭된 숫자와 함께 각각 매칭 변수에 관한 평균과 표준편차를 알 수 있다. 결과를 보면 33명의 대조군 중 17명과 실험군에서 17명이 매칭되었고 대조군에서의 16명이 unmatched 되었다(Fig. 5).

매칭 후의 성향 점수 분포에 대한 다양한 plot들을 아래와 같이 그려볼 수 있다.

```
>plot(matchingdata, type = "hist")
```

```
>plot(matchingdata, type = "jitter")
```

Jitter plot의 경우에 위의 명령어를 입력하고 엔터키를 친 후, 그래프로 가서 하나의 값을 클릭하고 ESC 버튼을 누르면 그 값을 확인할 수 있도록 되어 있다. 새로운 명령어를 입력하기 전에 확인하고 싶은 값이 없는 경우에는 바로 ESC 버튼을 눌러 다음 명령어가 작성될 수 있도록 하는 과정이 필요하다(Fig. 6).

예를 들어 확인하고 싶은 값을 클릭한 후 ESC 버튼을 누르면 값이 그래프에도 보여지고 입력창에도 보여지는 것을 알 수 있다(Fig. 7).

Histogram과 Jitter plot의 control group을 보면 매칭한 후 보정된 것을 볼 수 있다(Fig8, 9)⁸⁾.

이렇게 매칭된 데이터는 아래의 명령어를 통해 불러올 수 있고 17쌍으로 매칭되었기 때문에 34명의 데이터가 보여진다. 첫 번째 명령어는 매칭된 데이터를 matchingdata로 정렬시키고 두 번째 명령어는 이 데이터를 그림 9와 같이 보여준다. 매칭된 결과를 보여주는 subclass라는 변수가 1~17까지 생성된 것을 알 수 있다(Fig. 10).

```
>matchdata = match.data(matchingdata)
>matchdata
```

매칭된 결과를 다시 엑셀 file로 저장하는 과정으로 저장할 경로를 ‘ ’ 안에 입력하면 .CSV 확장명으로 저장할 수 있다.

```
>write.csv(matchdata, file = 'c:/PSM/
  matchdata.csv')
```

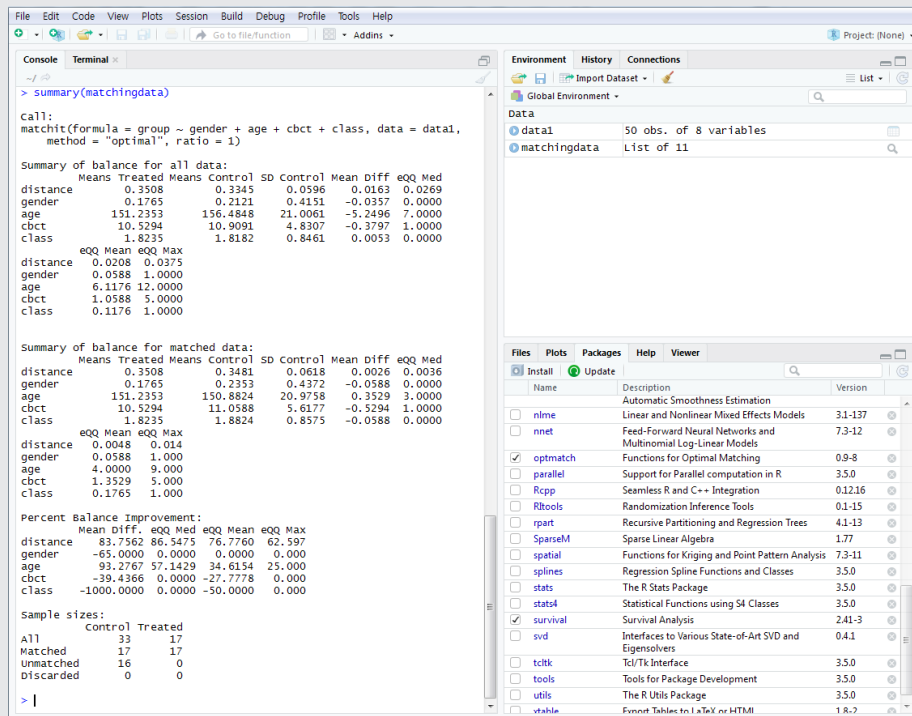


Fig. 5. summary 명령어를 통해 보이는 통계수치

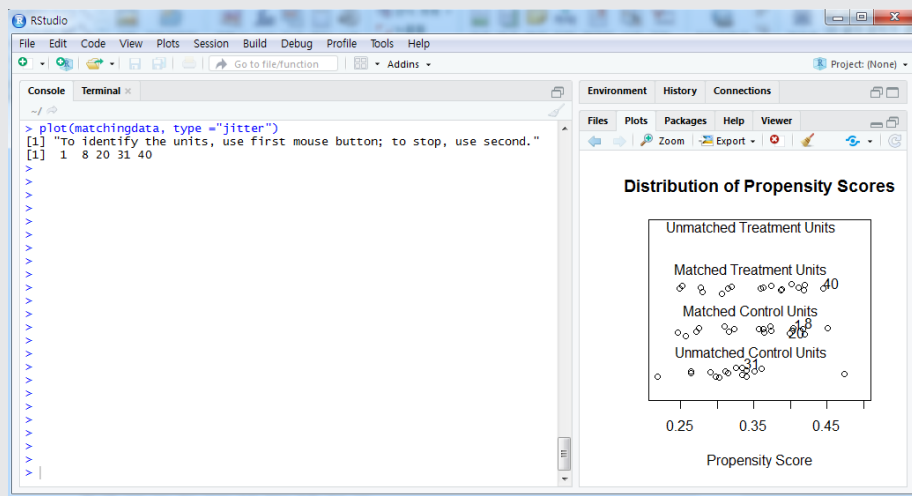


Fig. 6. Jitter plot 명령어를 입력하였을 때 뜨는 설명

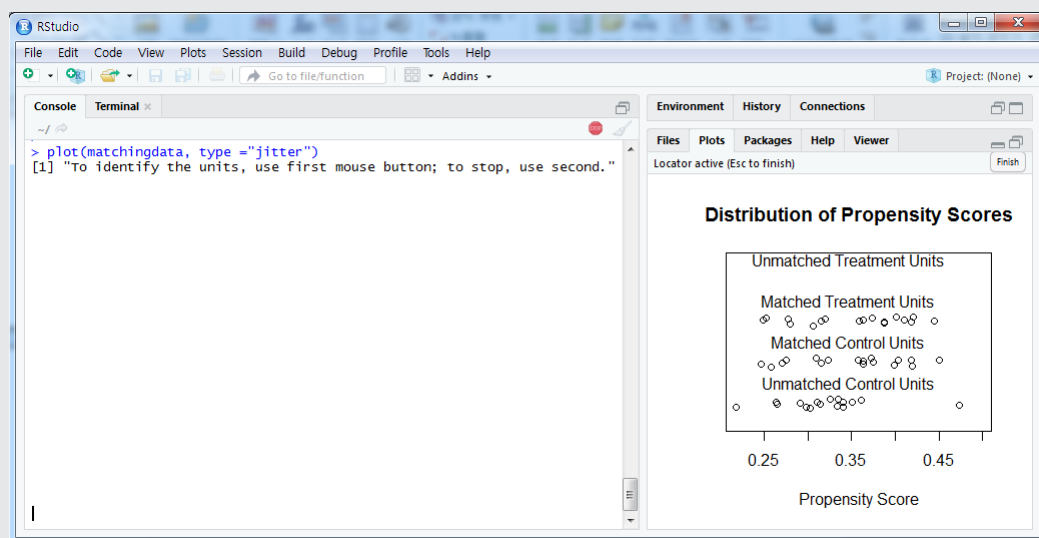


Fig. 7. Jitter plot에서의 특정한 값 확인



Fig. 8. Histogram

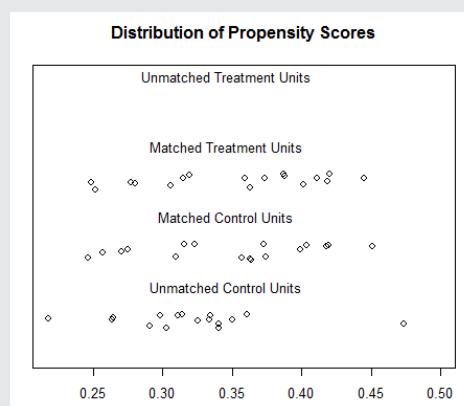


Fig. 9. Jitter plot

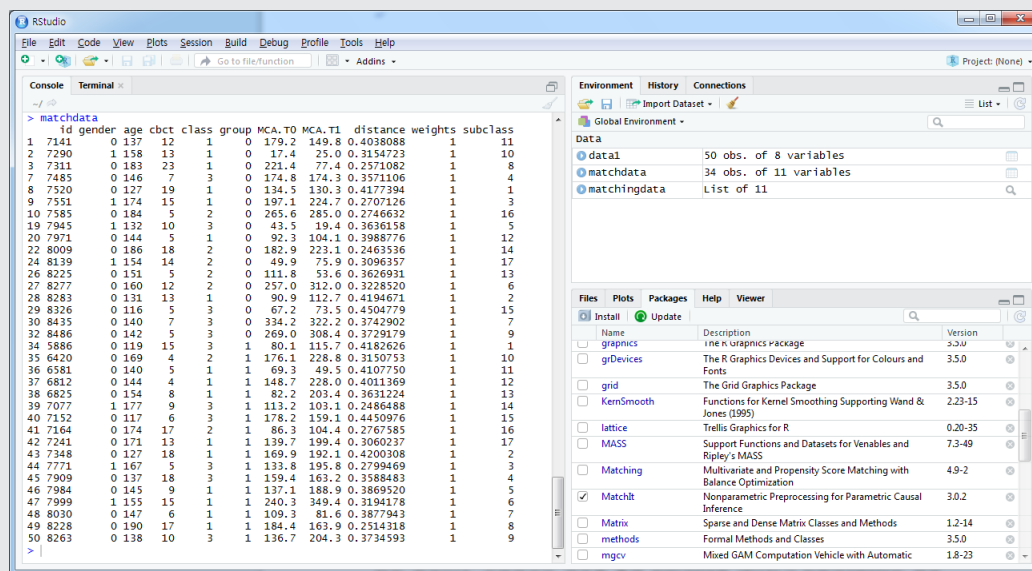


Fig. 10. 매칭 후의 데이터를 불러오는 과정

IV. 분석 과정

성향 점수를 이용하여 매칭시킨 실험군과 대조군 사이의 교정 전후의 MCA(Minimum Cross sectional Area)에 차이가 있는지를 알아보기 위해 분석을 하는 과정이다. 지금까지의 과정을 통해 matching된 데이터가 matchdata라는 파일명으로 저장된 것을 보여준다. 17쌍으로 매칭되었기 때문에 34명으로 구성된 데이터가 아래와 같다(Fig. 11).

```
>data2 = read.csv(file.choose())
```

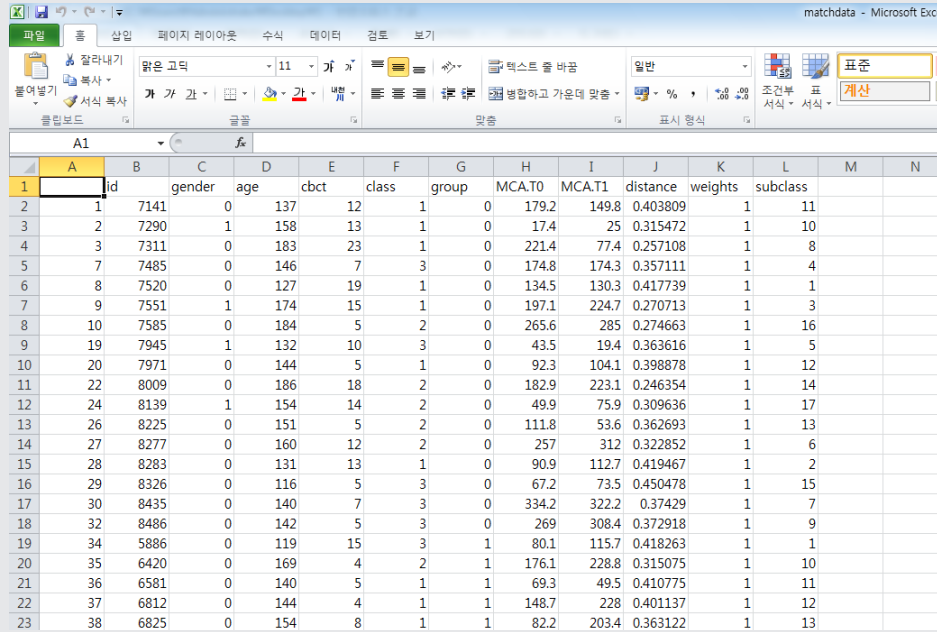
아래의 명령어를 이용하여 파일을 불러온다(Fig. 12). Data2가 제대로 불러왔는지 확인하고 다른 장문의 명령어 없이 변수명으로 데이터에 대한 정보를 불러올 수

있도록 데이터를 attach하는 과정이 필요하다. data2를 attach하기 전에 변수들의 중복을 피하기 위하여 data1을 detach해야 할 필요가 있다.

```
>detach(data1)
>data2
>attach(data2)
```

먼저 matching된 데이터의 교정 전 후의 MCA 차이 값(MCA.T1-MCA.T0)을 계산하여 MCA라는 변수로 만들고 data2 파일에 추가하여 data3라는 이름으로 저장한다. MCA라는 변수가 생성된 것을 알 수 있다(Fig. 13).

```
>MCA<-MCA.T1-MCA.T0
>data3<-cbind(data2,MCA)
```



	A	B	C	D	E	F	G	H	I	J	K	L	M	N
1		id	gender	age	cbct	class	group	MCA.T0	MCA.T1	distance	weights	subclass		
2	1	7141	0	137	12	1	0	179.2	149.8	0.403809	1	11		
3	2	7290	1	158	13	1	0	17.4	25	0.315472	1	10		
4	3	7311	0	183	23	1	0	221.4	77.4	0.257108	1	8		
5	7	7485	0	146	7	3	0	174.8	174.3	0.357111	1	4		
6	8	7520	0	127	19	1	0	134.5	130.3	0.417739	1	1		
7	9	7551	1	174	15	1	0	197.1	224.7	0.270713	1	3		
8	10	7585	0	184	5	2	0	265.6	285	0.274663	1	16		
9	19	7945	1	132	10	3	0	43.5	19.4	0.363616	1	5		
10	20	7971	0	144	5	1	0	92.3	104.1	0.398878	1	12		
11	22	8009	0	186	18	2	0	182.9	223.1	0.246354	1	14		
12	24	8139	1	154	14	2	0	49.9	75.9	0.309636	1	17		
13	26	8225	0	151	5	2	0	111.8	53.6	0.362693	1	13		
14	27	8277	0	160	12	2	0	257	312	0.322852	1	6		
15	28	8283	0	131	13	1	0	90.9	112.7	0.419467	1	2		
16	29	8326	0	116	5	3	0	67.2	73.5	0.450478	1	15		
17	30	8435	0	140	7	3	0	334.2	322.2	0.37429	1	7		
18	32	8486	0	142	5	3	0	269	308.4	0.372918	1	9		
19	34	5886	0	119	15	3	1	80.1	115.7	0.418263	1	1		
20	35	6420	0	169	4	2	1	176.1	228.8	0.315075	1	10		
21	36	6581	0	140	5	1	1	69.3	49.5	0.410775	1	11		
22	37	6812	0	144	4	1	1	148.7	228	0.401137	1	12		
23	38	6825	0	154	8	1	1	82.2	203.4	0.363122	1	13		

Fig. 11. 매칭 된 데이터

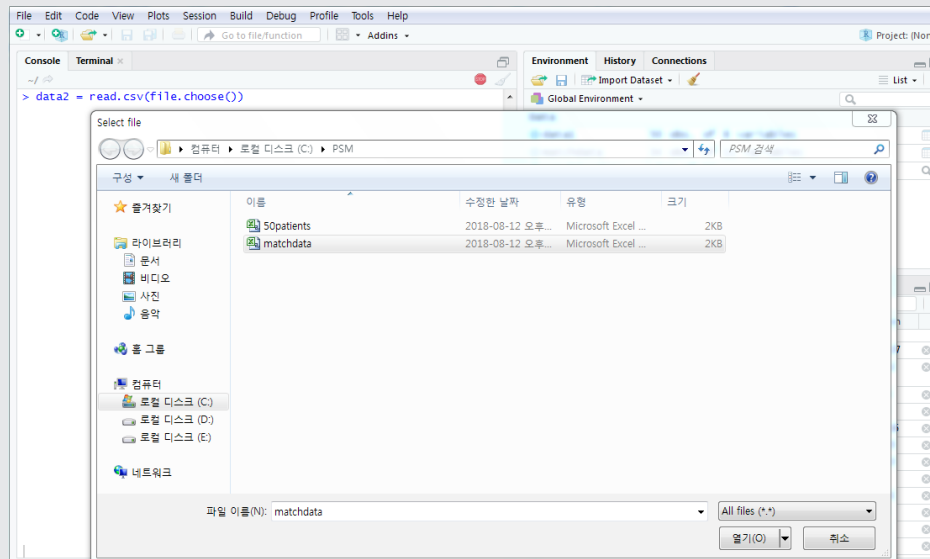


Fig. 12. 매칭된 데이터를 불러오는 과정

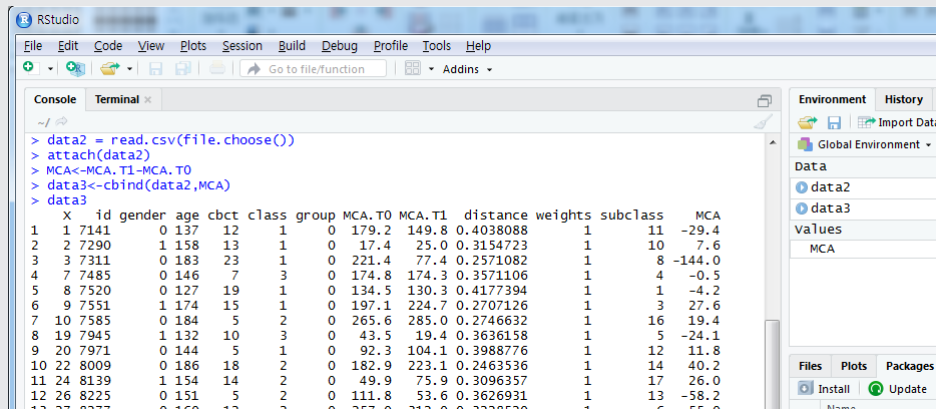


Fig. 13. MCA 변수를 기존의 데이터에 추가하는 과정

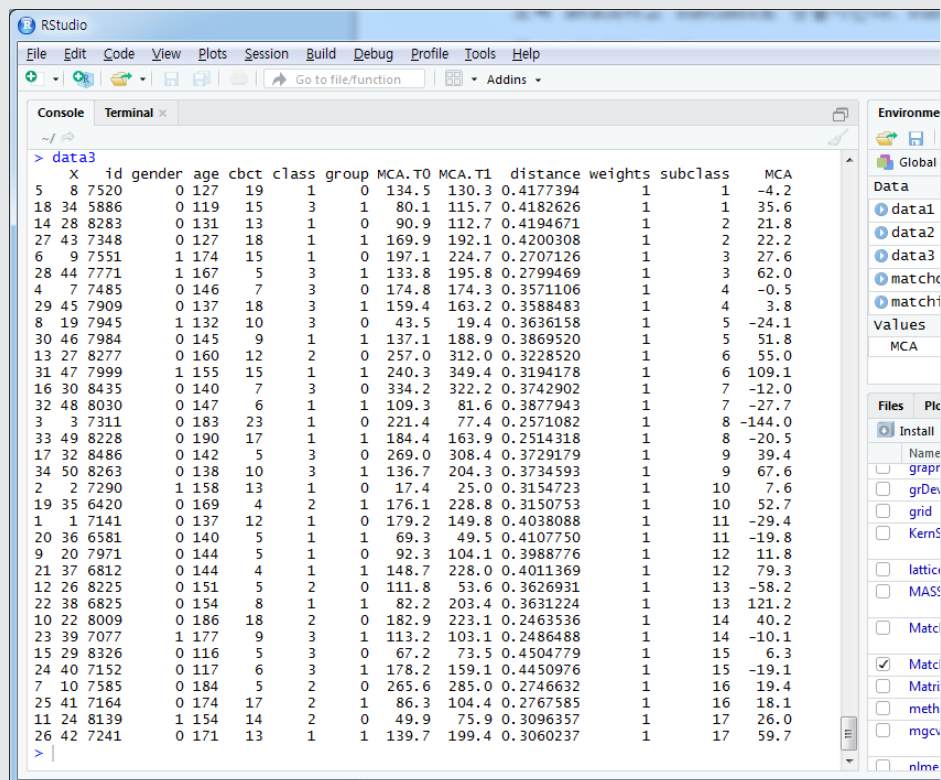


Fig. 14. paired t test를 위해 정렬된 데이터

```
>detach(data2)
>data3<-data3[c(order(data3$subclass)),]
>attach(data3)
```

Data3를 다른 장문의 명령어 없이 변수명으로 데이터에 대한 정보를 불러올 수 있도록 attach하고 subclass로 정렬시킨다. subclass가 1부터 순서대로 나열된 것을 볼 수 있다(Fig. 14). 매칭된 데이터끼리 분석해야 하므로 반드시 데이터 정렬한 후, group간의 paired t-test

를 시행한다(Fig. 15).

```
>t.test(MCA ~ group, data=data3,
paired=TRUE)
```

분석결과를 보면 $p=0.01823$ 로 0.05보다 작으므로 귀무가설을 기각한다. 따라서 case group과 control group사이의 교정 전후의 MCA값에 차이가 있다고 할 수 있다.

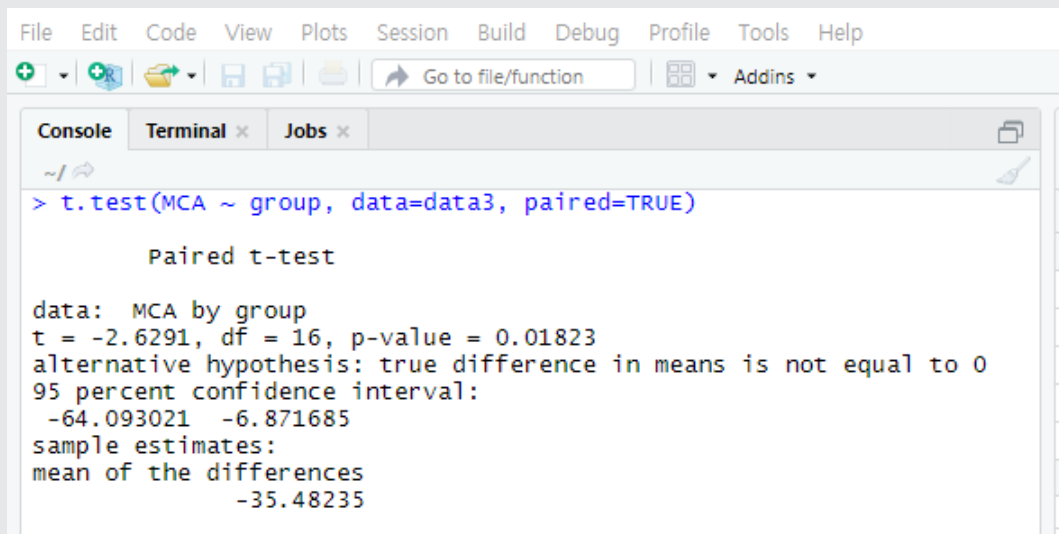


Fig. 15. paired t test 결과

V. 논의

기존의 1:1 매칭과 성향점수매칭의 비교

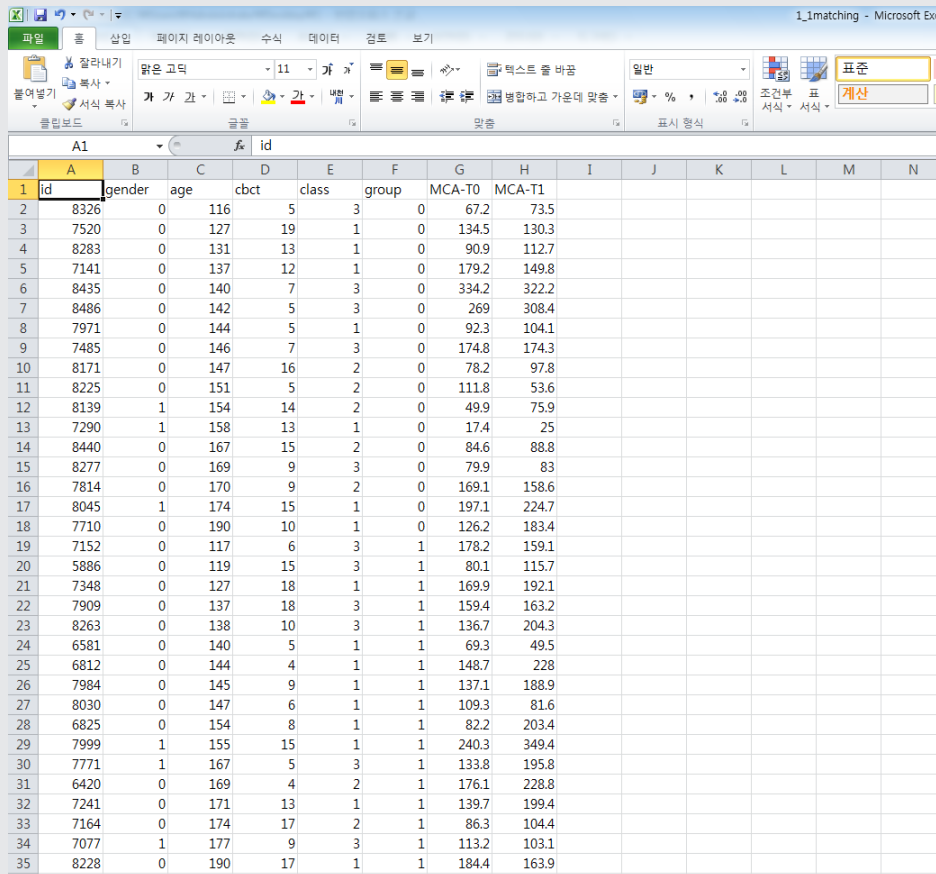
성향점수매칭방법을 설명하기 위해 사용된 실험군 17명과 대조군 33명의 가상의 데이터를 기존의 1:1 matching 방법인 나이와 성별에 따라 매칭한 데이터이다. 이 데이터도 전남대학교 치의학전문대학원 (<http://dent.jnu.ac.kr>) 일반대학원 자료실 게시판 번호 30번 Propensity score matching data를 클릭하여 1_1matching 데이터를 다운 받을 수 있다.

통상적인 연구에서 실험군과 대조군을 매칭하는 경우에 성별과 나이대가 동일하도록 단순 1:1 매칭을 시행하였다. 우리가 사용한 50명의 환자의 데이터를 나이와 성별에 대한 단순 1:1 매칭한 결과와 나이와 성별뿐만 아니라 CBCT 촬영 간격, class 관계까지 고려한 PSM 결과를 비교해봄으로써 PSM 방법의 필요성을 더욱 알 수 있을 것이다. 50명의 환자에 대하여 나이와 성별에 대하여 1:1 매칭한 데이터에 대하여 paired t-test를 시행한 결과 $p=0.08941$ 로 0.05보다 크므로 귀무가설을 기각하지 않는다. 따라서 실험군과 대조군 사이의 교정 전 후의 MCA값에 차이가 없다고 할 수 있다. 이 결과는 PSM의 결과와 차이가 있다. 치료결과에 영향을 주는 공변량인 CBCT와 class관계까지 고려한 PSM의 방법이 더 정확할 것이라고 할 수 있다. 따라서 특히 관찰연구를 함에 있어 치료의 효과를 비교하고자 할 때 결과에 영향을 미

치는 공변량의 선택과 함께 올바른 통계 방법의 선택이 올바른 연구 결과 도출과 치의학 발전에 기여할 수 있을 것이라 사료된다.

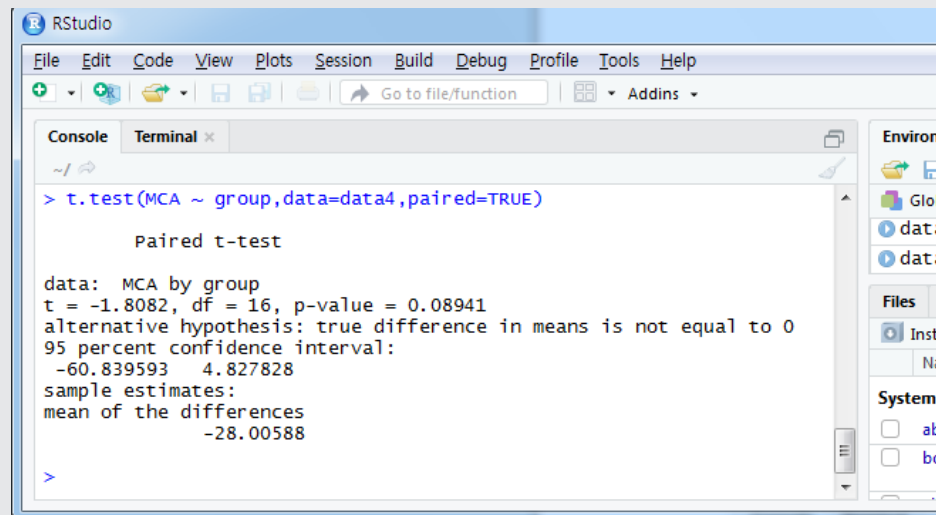
VI. 결 론

본 논문에서는 치의학 연구에 적합한 통계 방법들 중 하나로 성향점수 매칭 방법의 소개와 함께 무료 배포 프로그램인 R-studio를 사용하여 예제를 통한 PSM의 적용과정을 단계별로 설명하였다. 치의학 영역의 연구는 관찰연구가 대부분이며 처치를 통한 치료 효과를 알아보고자 하는 연구가 많고 특히 의학 연구에 비하여 표본의 수가 적은 것이 특징이다. 관찰연구의 제한점은 임의화의 부족이며 기저 특성의 차이에 의해 임의화를 완전히 시행하기란 어렵다. 그렇기 때문에 PSM 방법에서의 성향점수를 통해 공변량에 대한 편향을 보정하고 매칭함으로써 관찰연구를 강화할 수 있는 대안이 될 수 있을 것이라 생각된다. 특히 표본수가 적을 때, 공변량당 10개 미만이거나 공변량보다도 표본이 적은 경우에 model을 설정할 수 없기 때문에 PSM의 사용은 통계적으로 유리하며 필요하다고 할 수 있다^{9,10}. 성향점수분석 방법은 주어진 조건을 완벽하게 매칭한 것은 아니지만 성향점수를 통해 공변량에 대한 편향을 보정하고 매칭함으로써 관찰연구의 단점을 보완할 수 있는 대안이 될 수 있을 것이라 생각된다.



	A	B	C	D	E	F	G	H	I	J	K	L	M	N
	id	gender	age	cbct	class	group	MCA-T0	MCA-T1						
2	8326	0	116	5	3	0	67.2	73.5						
3	7520	0	127	19	1	0	134.5	130.3						
4	8283	0	131	13	1	0	90.9	112.7						
5	7141	0	137	12	1	0	179.2	149.8						
6	8435	0	140	7	3	0	334.2	322.2						
7	8486	0	142	5	3	0	269	308.4						
8	7971	0	144	5	1	0	92.3	104.1						
9	7485	0	146	7	3	0	174.8	174.3						
10	8171	0	147	16	2	0	78.2	97.8						
11	8225	0	151	5	2	0	111.8	53.6						
12	8139	1	154	14	2	0	49.9	75.9						
13	7290	1	158	13	1	0	17.4	25						
14	8440	0	167	15	2	0	84.6	88.8						
15	8277	0	169	9	3	0	79.9	83						
16	7814	0	170	9	2	0	169.1	158.6						
17	8045	1	174	15	1	0	197.1	224.7						
18	7710	0	190	10	1	0	126.2	183.4						
19	7152	0	117	6	3	1	178.2	159.1						
20	5886	0	119	15	3	1	80.1	115.7						
21	7348	0	127	18	1	1	169.9	192.1						
22	7909	0	137	18	3	1	159.4	163.2						
23	8263	0	138	10	3	1	136.7	204.3						
24	6581	0	140	5	1	1	69.3	49.5						
25	6812	0	144	4	1	1	148.7	228						
26	7984	0	145	9	1	1	137.1	188.9						
27	8030	0	147	6	1	1	109.3	81.6						
28	6825	0	154	8	1	1	82.2	203.4						
29	7999	1	155	15	1	1	240.3	349.4						
30	7771	1	167	5	3	1	133.8	195.8						
31	6420	0	169	4	2	1	176.1	228.8						
32	7241	0	171	13	1	1	139.7	199.4						
33	7164	0	174	17	2	1	86.3	104.4						
34	7077	1	177	9	3	1	113.2	103.1						
35	8228	0	190	17	1	1	184.4	163.9						

Fig. 16. 기존의 1:1 matching 과 성향점수 매칭의 비교를 위해 사용된 데이터



```

RStudio
File Edit Code View Plots Session Build Debug Profile Tools Help
~ / Go to file/function Addins
Console Terminal x
> t.test(MCA ~ group, data=data4, paired=TRUE)

Paired t-test

data: MCA by group
t = -1.8082, df = 16, p-value = 0.08941
alternative hypothesis: true difference in means is not equal to 0
95 percent confidence interval:
-60.839593 4.827828
sample estimates:
mean of the differences
-28.00588
>

```

Fig. 17. 기존의 1:1 matching의 paired t test 결과

Conventional 1:1 matching

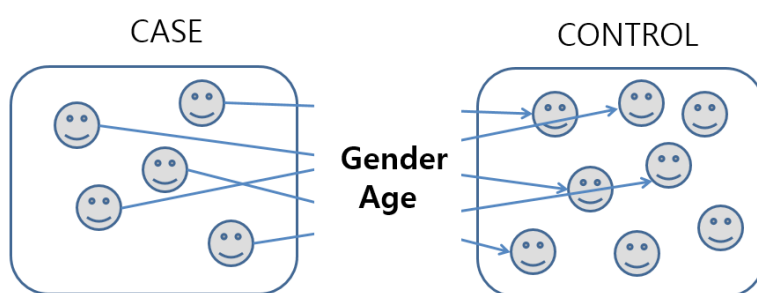


Fig. 18. Conventional 1:1 matching 다이어그램

Propensity score matching

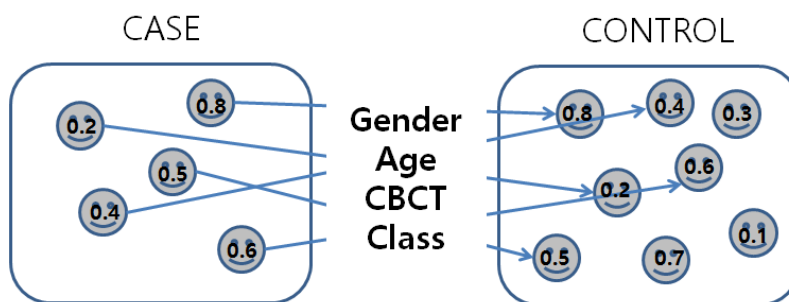


Fig. 19. Propensity score matching 다이어그램

• 참고 문헌 •

1. 니시후치 시로무. 빅데이터를 지배하는 통계의 힘. 1판. 서울. 비전비엔 피. 2013.
2. Park SH, Lim HJ. A step-by-step guide to Meta-analysis with dichotomous outcomes using RevMan in dental research. The Journal of the Korean dental association 2018;56(1):18-40.
3. Lim HJ, Park SH. A step-by-step guide to Generalized Estimating Equations using SPSS in dental research. The Journal of the Korean dental association 2016;54(11):850-864.
4. Lim HJ. Sample size determination in dental research. The Journal of the Korean dental association 2014;52(9):558-569.
5. Lim HJ. Meta-analysis in dental research. The Journal of the Korean dental association 2014;52(8):478-490.
6. Olmos A, Govindasamy P. Propensity Scores: A Practical Introduction Using R. J Multidiscip Eval. 2015;11(25):68-88.
7. Randolph JJ, Falbe K, Manuel AK, Balloun JL. A Step-by-Step Guide to Propensity Score Matching in R. Practical Assessment Research&Evaluation. 2014;19(18):1-6.
8. Lee DK. An introduction to propensity score matching methods. Anesth Pain Med. 2016;11(2):130-148.
9. Austin PC. An Introduction to Propensity Score Methods for Reducing the Effects of Confounding in Observational Studies. Multivariate Behav Res. 2011;46(3):399-424.
10. Sainani KL. Propensity scores: uses and limitations. Phys Med Rehabil. 2012;4(9):693-697.