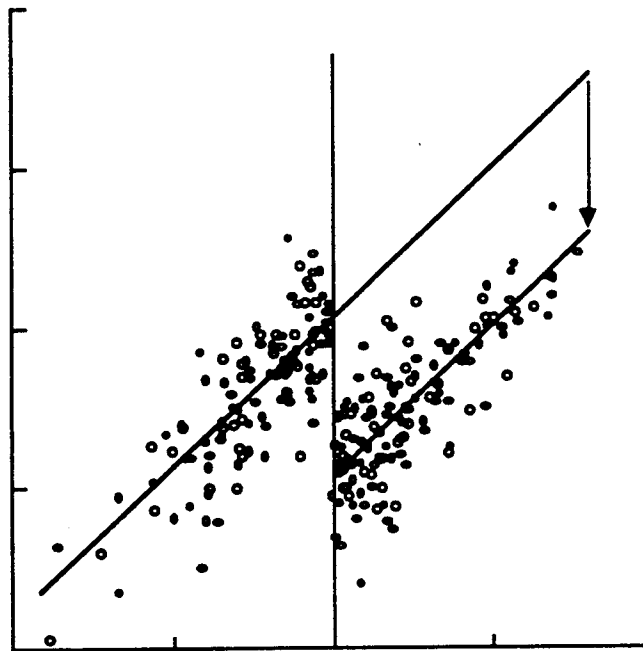


The Regression - Discontinuity Design: An Introduction

Research Methods Paper Series
Number 1

William M.K. Trochim
Cornell University



This paper is part of a series. To obtain a copy of this paper or for information about other papers in this series, please contact:

Research Methods Paper Series
Thresholds National Research and Training Center on Rehabilitation and Mental Illness
2001 N. Clybourn Avenue
Suite 302
Chicago IL 60614
(312) 348-5522

The Regression - Discontinuity Design: An Introduction

This is an introduction to the regression-discontinuity design and how it works. The regression-discontinuity (from here on we'll use the label RD) design is not well known and has received little use in mental health, rehabilitation, mental illness, and related areas. But it has great potential for studying the effects of treatments in these areas. The major advantage of the RD design is that, unlike some other common designs, it allows you to give the treatment you are studying to the people who most need or deserve it. Some designs require that, in order to conduct a good scientific test, you deny the treatment from some people who might need it. These people act as comparison cases for the people who receive the treatment and show what happens when treatment is not given. While denying the treatment to some people for the sake of a scientific test might be good science, it is often in conflict with good clinical practice. The clinician wants to help an ill person -- not deny them treatment. The RD design allows you to give the treatment to those who are most impaired and, if carried out correctly, is also a credible scientific method for studying the effects of the treatment. The real potential of the RD design is that it addresses both good scientific and good clinical practice. Although the name of the design sounds technical, it is actually not too difficult a strategy to understand. This paper outlines the RD design, describes how it works, introduces the role of regression analysis, and shows how the RD design is distinguished from other similar designs.

What is a Research Design?

A *research design* describes how the sample, measures and treatments are put together over time when you conduct a study to assess the effects or outcomes of a program or treatment. Every design has certain advantages and disadvantages. For instance, some designs require *control groups* -- groups that don't receive any treatment and are used to compare against treated persons or program recipients. The advantage of control groups is that they enable the researcher to see what would have happened to the treated people if they had not been given treatment. But control group designs usually require more people in your study. This tends to increase both the study costs and the length of time needed to complete the study. Other designs require that you measure the same group of people on several occasions, such as before and after a program or treatment. The advantage of repeated measurement is that one is able to look at change over time especially from before to after the treatment or program. Again, however, this will be more costly and time consuming than designs which only measure at one occasion. When you select a research design, you have to balance trade-offs like these. You have to weigh how much scientific credibility and accuracy you want against the costs and negative qualities associated with the design.

What is the Regression-Discontinuity Design?

The RD Design is a *before-and-after two group design* (sometimes also called a pre-post two group design). This means that all persons in the study are measured both before and after the treatment or program. In addition, there are always two groups, usually one that receives the treatment and one that does not. But there are many other designs which can also be called before-and-after two group designs. What makes the RD design unique -- and sets it apart from all other designs -- is the in the RD design persons are assigned to groups solely based on a cutoff score on a preprogram measure.

Let's consider a simple example of an RD design. Imagine that you wanted to conduct an outcome evaluation of the effectiveness of a new program for treating mental illness. The first thing you would need is a "before" measure, or a *pretest*. Let's imagine that we have a test which measures the severity of the person's mental illness. High scores would indicate that the person is more severely ill. Our new program has fewer patient slots than available patients and so we will have to

Person	X	Y
1	52.503	52.759
2	47.579	47.777
3	53.245	53.76
4	50.334	50.923
5	49.511	49.841
6	50.523	50.837
7	50.592	52.265
8	52.531	53.283
9	48.643	47.767
10	50.407	52.471
11	57.003	58.693
12	47.04	45.692
13	52.698	55.738
14	48.707	47.597
15	46.899	47.488
16	46.85	47.854
17	50.513	52.596
18	49.425	48.437
19	51.896	53.066
20	48.75	49.063
21	51.782	48.879
22	49.871	50.121
23	51.042	50.412
24	46.586	46.824
25	51.813	55.168
26	45.45	45.882
27	55.532	53.404
28	48.445	49.359
29	49.479	51.734
30	51.296	50.317
.		
.		
.		
250	48.891	48.875

Table 1. Pretest and one-month posttest scores for some of the 250 persons who present themselves to our agency.

decide which patients get the program. Let's assume that we want to target the program to those patients who will need it most. Here, we might think that those who are most severely ill will need the most help, and so they are the ones we would most like to receive our treatment. To construct the *assignment criterion*, let's say we consult with clinicians, caseworkers, and other relevant groups to get them to decide on a *cutoff score* on our measure of severity of illness. All those who score at or above this cutoff score on our measure will be assigned to our new program, and will be called the *program group* or *treatment group*. Those who score below the cutoff will not be eligible for the treatment, and will be considered *control group cases*. After the program has been given, we can measure its outcome on one or more variables. Often, we want to measure the same thing we measured on the pretest, in this case, severity of illness. But we can measure other outcomes as well, such as coping ability, interpersonal skills, job-related skills, recidivism, and so on. Each of these is an "after" measure or *posttest*.

All persons in the study -- those who scored above and below the cutoff score on the pretest -- must be measured on the posttests. This example illustrates a simple RD design. Notice that the key characteristic that tells you this is an RD design is that a cutoff score is used to determine whether persons get the program or not.

How Does the RD Design Work?

An Imaginary Mental Health Evaluation. Let's begin by imagining two persons who come into our agency and who wish to take part in our new program for treating mental illness. We measure each of them on our test of severity of illness. The first person gets a pretest score of 52.503 while the second scores lower with 47.579. Of course, there are many more people who come to the agency than just these two. Let's imagine that over a period of time, say, one year, a total of 250 people come to the agency seeking help. Furthermore, let's imagine that no treatment is actually given to any of these patients for four weeks (maybe we have a

lengthy waiting list!). At the end of the four week period we will measure each person (of the 250) on the posttest measure of severity of illness. Let's look at what the database of the pre and posttest scores might look like for a subsample of the 250 persons.

In Table 1 we see the pretest and posttest scores for the first thirty persons who came to the agency. The scores for the first two people are shown in boldface type. We will follow typical statistical conventions and label the pretest with the letter 'X' and the posttest with the letter 'Y.' The test scores might be average values across a large number of different test items, where the overall average of the test is expected to be 50 for this type of population. Notice that the pretest and posttest scores seem to hover around 50 and that for most people there is not a great difference between their pretest and posttest. That is, we assume that people don't change much with respect to severity of illness over a four-week period if they are left untreated. The slight difference between a person's pretest and posttest score can be attributed to the random occurrences which

affected testing on each occasion, such as disturbances in the testing setting, the client's mood and attentiveness on that day, the weather, and so on.

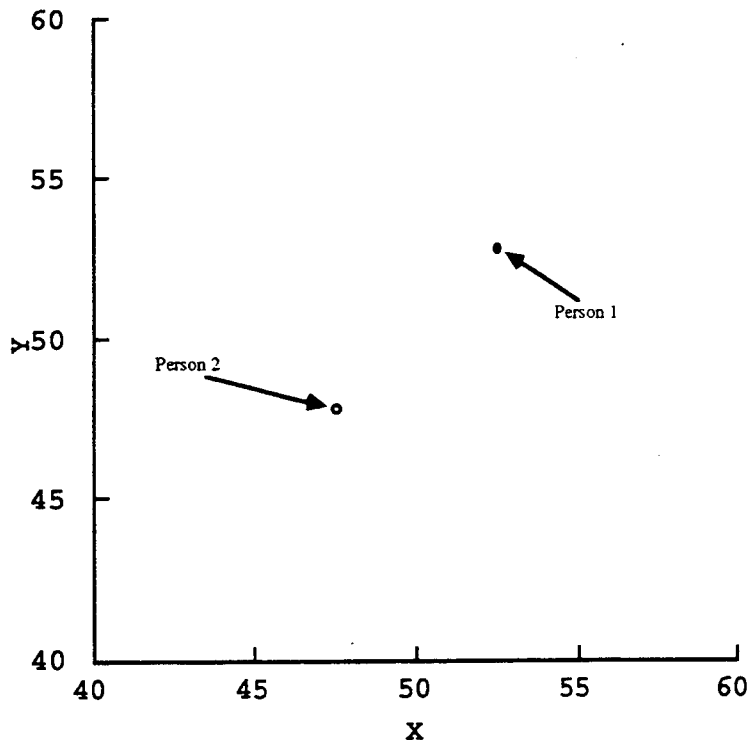


Figure 1. Bivariate plot showing the pretest and posttest scores for two persons.

We might also wish to view the first two persons using a graph. One way to do this is to graph their pretest score against their posttest score. This is called a *bivariate plot* because two (bi-) variables are plotted against each other. We usually put the posttest, Y, on the vertical axis and pretest, X, on the horizontal. The bivariate plot for the first two persons is shown in Figure 1. Each person has a single dot on the map which shows their pretest and posttest score. For instance, to see what the pretest score for Person 1 is, all you have to do is draw a line from that person's dot straight down to the X-axis. This would confirm that Person 1 scored a little bit higher than 52 on the pretest. To find out the posttest score for Person 1, all you have to do is draw a line from person 1's dot left across to the Y-axis.

This would verify that Person 1 scored at nearly 53 on the posttest. While Person 1 scored higher than Person 2 on the pretest, Person 1 also scored higher on the posttest. This means that Person 1 had a higher severity of illness on both tests. Person 2 scored a little lower on both the pretest and posttest than Person 1 did.

The bivariate plot shows the pretest and posttest scores for only the first two persons. Usually, we would want to graph the bivariate plot for the entire group of 250 persons. Such a plot is shown in Figure 2.

Notice that the "cloud" of 250 dots seems to move from the lower left to the upper right of the figure. This is the pattern we typically find when we have a *positive correlation* between the pretest and the posttest. This means that across the entire group of 250 people, persons who had higher pretest scores tended in general to have higher posttest scores as well. It also means that persons with lower pretest scores also tend to have lower posttest scores. We would have a *negative correlation* if the cloud moved from the upper left to the lower right. This would mean that high pretest scores are associated with low posttest ones and vice versa.

A bivariate plot of this type is a very convenient way to display data on two variables for many persons. Notice how much easier it is to see the positive correlation between the pretest and posttest on the graph than it would be to see it in a table of numbers like the one given in Table 1 earlier.

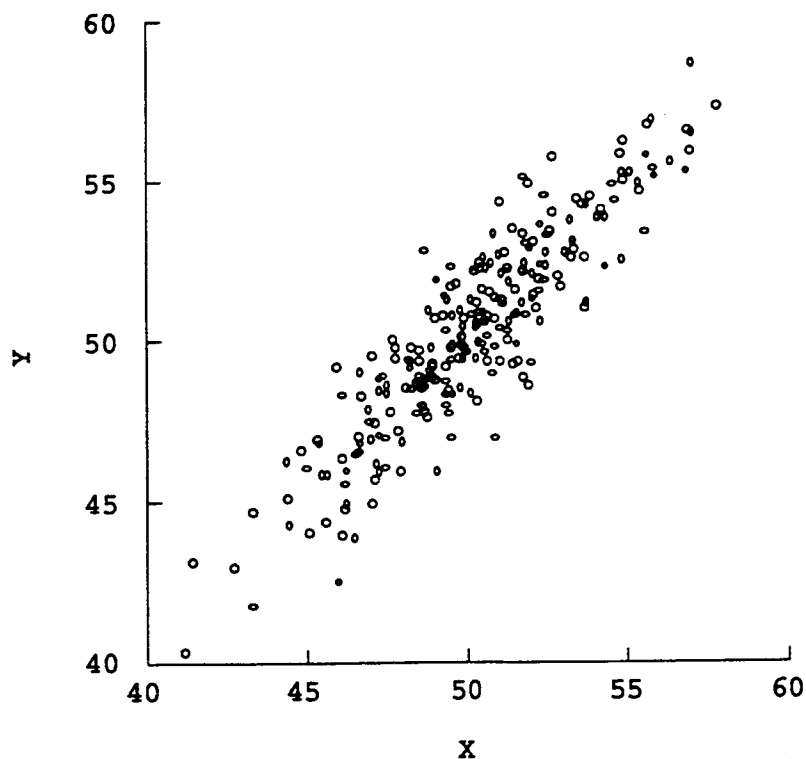


Figure 2. Bivariate plot showing the pretest and posttest scores for all 250 persons.

describes what would happen when nothing (no treatment) is given.

Now let's look at what would happen if we decided to use the RD design to study the effects of a treatment on some of these persons over the four-week period. Remember that with the RD design, persons are assigned to either receive the treatment or not using a cutoff point on a pretreatment measure. Let's imagine that before the study began we consulted with clinicians and staff members about who should get our new treatment. They all agreed that they would want the treatment to be targeted to those who are most severely ill.

Figure 2 shows the pre-post plot for the 250 persons, but where are the dots for Persons 1 and 2? Figure 3 shows where these two persons fall within the bivariate plot of the entire 250 people. Notice that these first two persons fall into the middle of the cloud of dots -- they are not "outliers" relative to the other 248 persons.

All of the plots shown so far assume that we just measure the 250 persons four weeks apart, but that they didn't receive any treatment during that four week period. The graphs therefore show what we would expect would happen to these 250 people if we don't give them a treatment over the four week period. Or, this is what a no-treatment control group plot would look like for the four week period. This is often called the *null case* graph because it

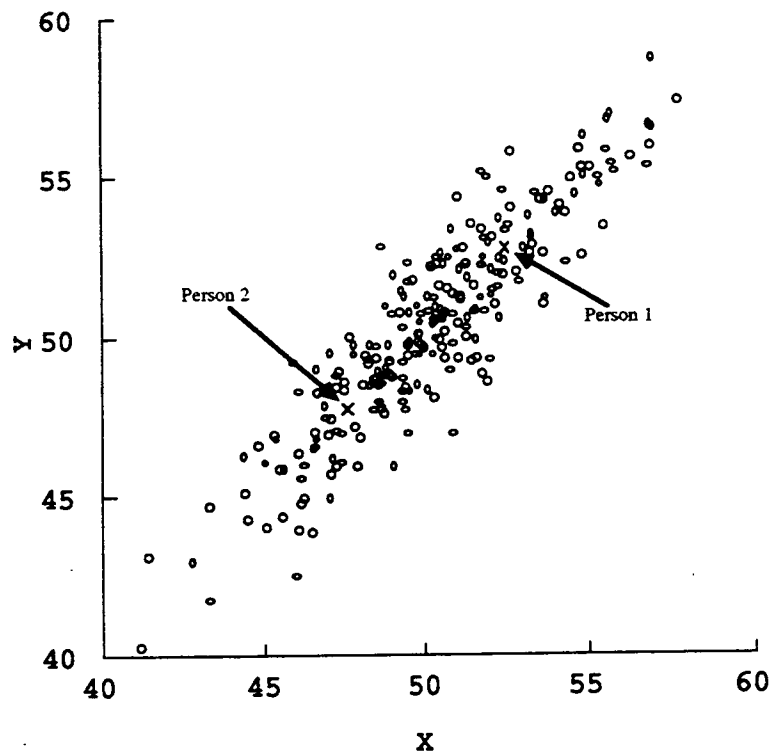


Figure 3. Bivariate plot showing pretest and posttest scores for all 250 persons with the first two persons highlighted.

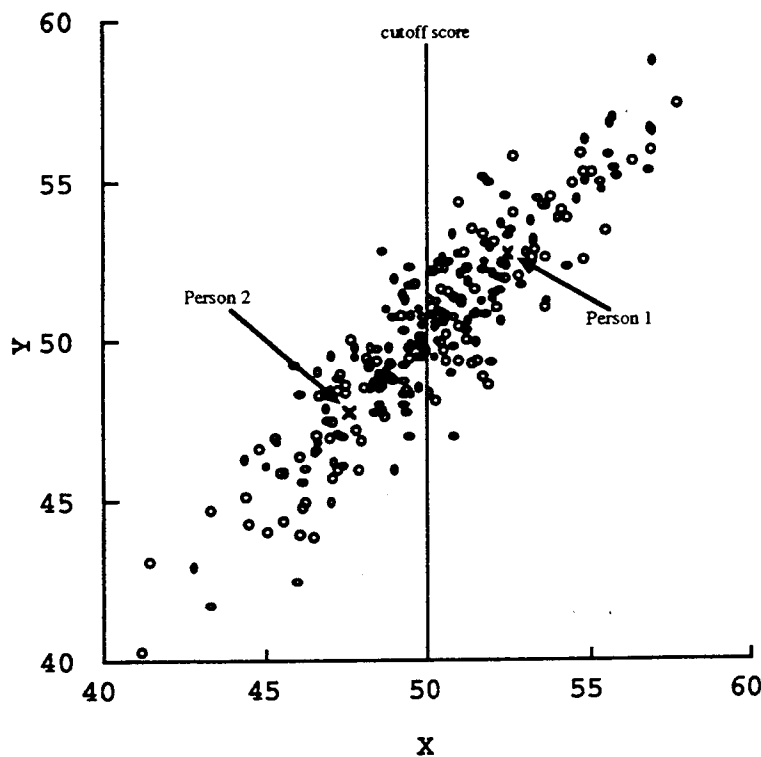


Figure 4. Bivariate plot with a cutoff score at the pretest value of 50 points.

Notice that according to the cutoff value, Person 1 would have a high enough pretest score to be eligible for the new treatment, whereas Person 2 would not.

Now, let's consider what would happen if our new treatment worked. Obviously, the treatment would not change any person's pretest score because those were measured before the program was given.

And, it should also be obvious that the posttest scores for all persons who scored below 50 on the pretest -- the control group -- would also not be affected, because they don't get any treatment. The only scores that would be affected would be the posttest scores of the persons in the treatment group -- those scoring at or above 50 points on the pretest. How would their

Based on past admissions rates and the number of treatment slots they have available, they agreed that a reasonable cutoff point could be the average of the pretreatment measure of severity of illness.

From past experience with that measure, they knew that for their typical patients the average score would be about 50 points. Therefore, they decided that they would use a cutoff value of $X=50$ to determine who would receive the treatment. This means that any person scoring at or above 50 points would automatically receive the new treatment while those scoring below 50 would be considered control cases and would not receive the new treatment. The bivariate plot for all 250 cases is shown with the cutoff value in Figure 4.

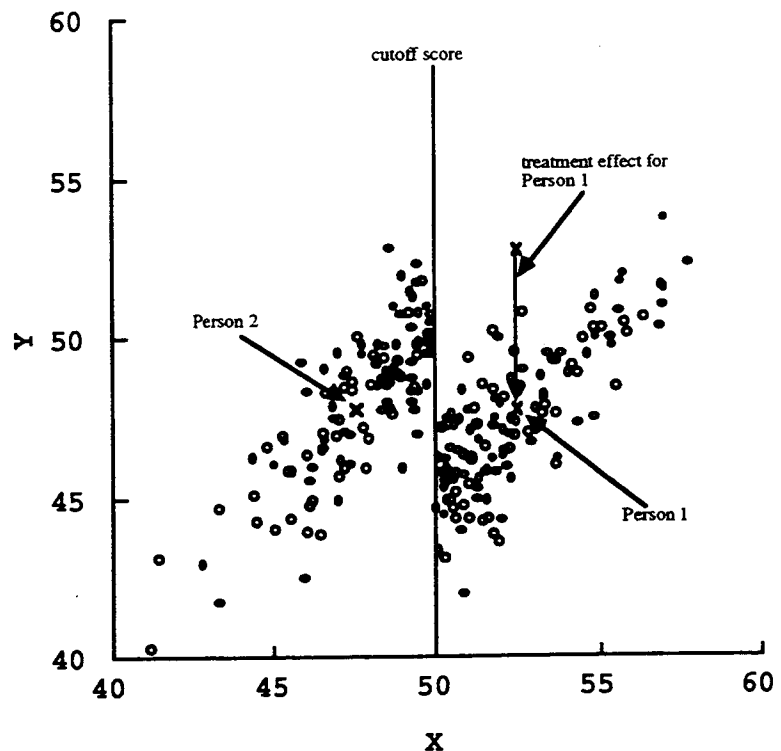


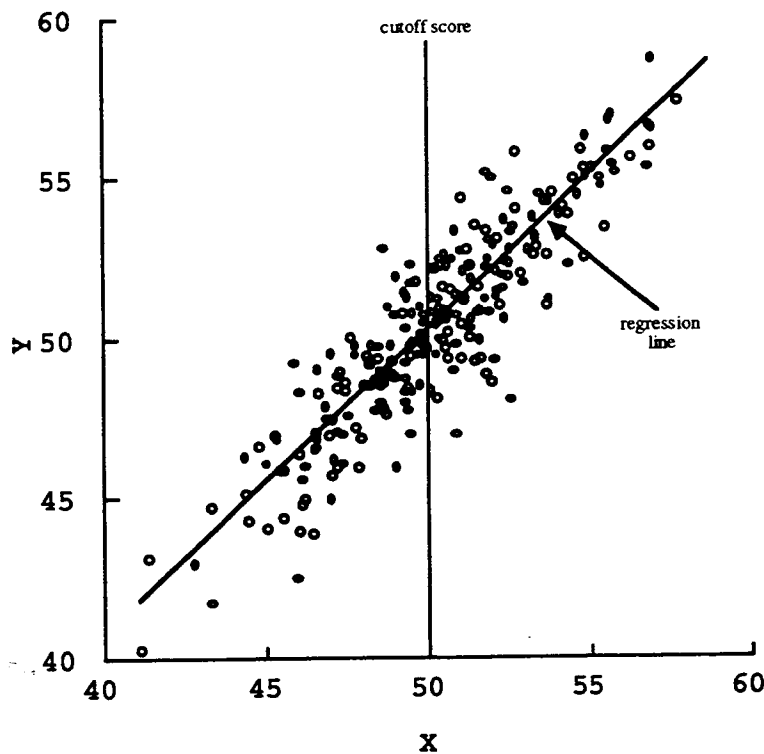
Figure 5. The effect of a treatment on the posttest scores of the treated persons.

scores be affected? If the treatment works, we would expect that it would make people less severely ill. This means that the treated persons would have lower posttest scores after being treated than they would have otherwise.

In Figure 5, we can see what the data would look like if the new treatment had a "positive" effect on all of the treated persons. Remember that a "positive" effect here means that the posttest severity of illness is lower. We can see in the figure what has happened for our first person. An arrow pointing straight down shows where Person 1 would have scored without the treatment and where that person actually did score when the treatment was given. If you measure the length of that arrow you will see that Person 1 had a posttest severity of illness score about 5 points lower than they would have had if they hadn't had the treatment. The same is true for all of the other persons in the treatment group -- their posttest scores are lower than we would have expected them to be without treatment.

The Role of Regression Analysis. There is only one problem with Figure 5. In real life, we can't have Person 1 receive and not receive the treatment at the same time! That is, we would never know for sure where both points which are joined by the arrow in Figure 5 would really be. After all, a person can't be in two places at once.

To help address this, we need to find some way to "guess" where Person 1 (and all the rest of the treated cases) would be if they didn't receive the treatment. To do this, we use *regression analysis*. Regression analysis is a way of describing relationships between variables statistically. In the RD design, we use regression analysis to describe the relationship between the pretest and the posttest measures. A regression analysis enables us to summarize a bivariate cloud of dots with a more simple shape, called a *regression line*.



Let's go back to the simple bivariate plot shown in Figure 4 and show what the straight-line regression line would look like for this plot. This is shown in Figure 6. Here, the regression line "slices" through the cloud of dots. That is, about half of the dots with the same pretest (X) score will fall above and below the regression line. Another way to look at this is like an average. The regression line is analogous to an average posttest value at each pretest value.

When we do a regression analysis, we always specify what shape line we would like to fit to the cloud of data. In this case, we fit a "straight" line. But there may be times when the cloud doesn't look like it would be well fit with a straight line and we might try some type of curved line instead. The regression line is often called "*the line of best fit*"

Figure 6. Bivariate plot with a straight regression line plotted through the cloud of points.

through the data. In Figure 6, it is clear that a straight line fits this cloud of data well.

We use the regression line as a simpler summary of the original data points. The regression line is also sometimes used for prediction. For instance, if we had the line in Figure 6, we could predict the posttest score for a person who has a specific pretest score. To do this, we would start at the pretest score value (on X) and draw a line straight up until we reached the regression line. Then, from that point on the regression line, we would draw a line horizontally straight to the left until we reached the posttest (Y) value. That value would be our predicted posttest score for the pretest score we started with.

Regression Analysis in the RD Design. In the RD design, we use regression analysis to "guess" where the treated cases would be if they had not received the treatment. This is shown in Figure 7. Figure 7 is exactly the same as Figure 5 except that the regression lines are also shown. This figure shows the basic RD design. Notice that the regression line that fits through the control group points is "projected" into the treatment group region. This projected line is our best guess of where the treated group cases would have fallen (on average) if they had not received the treatment or if the treatment had absolutely no effect. But in Figure 7, the treatment does have an effect -- the posttest scores for the treated person are lower than we would expect them to be. This is what is meant by the treatment effect in the RD design. The treatment effect is the degree to which the posttest values for the treated persons are different from what we expect. And, our expected treatment group performance is indicated by the projection of the control group line into the region of the treatment group. In Figure 7 we see that the *treatment effect* is equal to the length of the downward-pointing arrow. Notice that the treatment group has its own regression line.

If the treatment was never given the treatment or it had no effect at all, we would expect that the treatment and control group regression lines would be one continuous line as shown in Figure 6. But in Figure 7, we see that the regression lines for the two groups aren't one continuous line -- they differ by the amount shown in the downward arrow.

What Data Do We Need for a RD Design?

In order to conduct an RD design we minimally need three measures. First, we need a pretest or preprogram value for each person in the study. Second, we need a score which tells us which group (treatment or comparison) each person is in.

Usually we use a numerical value of '1' to indicate that the person is in the treatment group, and a value of '0' to show that they are in the control group. Finally, we need at least one posttest measure designed to show the effect of the treatment or program. For a subsample of the persons in our hypothetical study, the data for the RD design which is graphed in Figure 7 is shown in Table 2. In the table we follow convention by labeling the pretest with an 'X,' the group variable with a 'Z' and the posttest with a 'Y.'

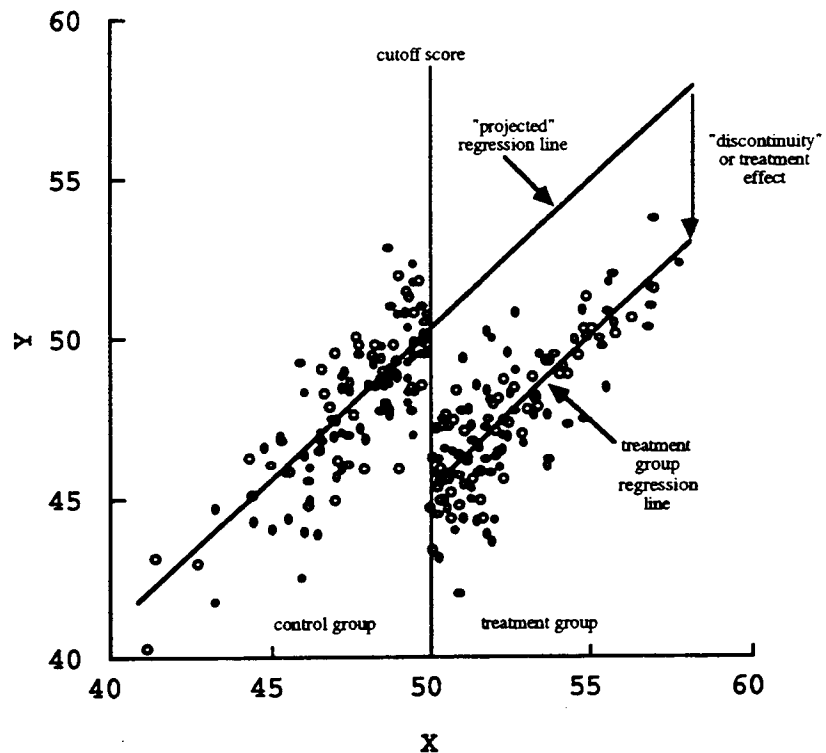


Figure 7. The bivariate plot for an RD design with regression lines shown.

Person	X	Z	Y
1	52.503	1	47.759
2	47.579	0	47.777
3	53.245	1	48.76
4	50.334	1	45.923
5	49.511	0	49.841
6	50.523	1	45.837
7	50.592	1	47.265
8	52.531	1	48.283
9	48.643	0	47.767
10	50.407	1	47.471
11	57.003	1	53.693
12	47.04	0	45.692
13	52.698	1	50.738
14	48.707	0	47.597
15	46.899	0	47.488
16	46.85	0	47.854
17	50.513	1	47.596
18	49.425	0	48.437
19	51.896	1	48.066
20	48.75	0	49.063
21	51.782	1	43.879
22	49.871	0	50.121
23	51.042	1	45.412
24	46.586	0	46.824
25	51.813	1	50.168
26	45.45	0	45.882
27	55.532	1	48.404
28	48.445	0	49.359
29	49.479	0	51.734
30	51.296	1	45.317
.			
.			
.			
250	48.891	0	48.875

Table 2. Pretest (X), treatment group variable (Z), and one month posttest scores (Y) for some of the 250 persons who present themselves to our agency (treatment effect assumed).

Notice the differences between this table and Table 1. Here, persons who have a pretest score below 50 (on X) are assigned to the control group ($Z=0$) while those at or above 50 are in the treatment group ($Z=1$). Also notice that all control group persons have the same posttest value in Tables 1 and 2. That's because they don't receive any treatment and therefore we wouldn't expect their posttest scores to be any different. The persons in the treatment group, on the other hand, show posttest scores which are 5 points lower than in Table 1. This is the treatment effect -- in this example, -5 points. The data in Table 2 can be entered into a regression analysis in almost any statistical computing package in order to conduct the RD analysis.

How Did the RD Design Get its Name?

We can now see how the RD design got its name. When there is a treatment effect, there is a discontinuity or jump in the regression lines at the cutoff point. Sometimes, it's convenient -- and clearer -- to show only the regression lines in the graph of an RD design. For the example given above, the regression lines are shown in Figure 8.

Here, we can more clearly see the discontinuity in the regression lines which the treatment seems to cause. Although the term "regression-discontinuity" is an accurate description, it is hardly a flattering name. In everyday language both parts of the name have connotations which are primarily negative. To most people "regression" implies a reversion backwards or a return to some earlier, more primitive state, while "discontinuity" suggests an unnatural jump or shift in what might otherwise be a smoother, more continuous process.

To a research methodologist, however, the term regression-discontinuity carries no such negative meaning. Instead, the RD design is seen as a useful method for determining whether a program or treatment is effective. Some other names have also been suggested for the same design. For instance, some people refer to the RD design as the "cutoff-based" research design to highlight its most unique feature. Others prefer to call it a "rule-based assignment" research design, although that term might actually lead to confusing RD with other rule-based designs. In any even, we use the traditional name because the design has been most often described that way in the research literature.

How is the RD Design Described in the Research Literature?

If you ever contemplate using the RD design you will at some point need to describe it in technical terms. The RD design can be described in a number of different ways. Some would call it a "pre-post two-group quasi-experimental design with assignment to group by cutoff." Others would substitute the term "experimental" for quasi-experimental in the description. The distinction between an experiment and a quasi-experiment is not clear, and depending on your definition of an experiment, the RD design may be either. For instance, some people (especially those coming from the medical research tradition)

argue that an experiment is any research design where assignment to group is controlled by the researcher. In this sense, the RD design is an experiment. Others (especially those who come from psychology, education, or program evaluation) define an experiment as any design which utilizes random assignment to group. Clearly, in this sense, the RD design would not qualify and would probably be classified a quasi-experiment.

Another way to describe a design is through *design notation*. For instance, the RD design might be notated:

C	O	X	O
C	O		O

where the C indicates that groups are assigned by means of a cutoff score, an O stands for the administration of a measure (pretest, posttest) to a group, an X depicts the administration of the program or treatment, time moves from left to right, and each group is described on a single line (i.e., program or treatment group on top, control or comparison group on the bottom). Different writers vary slightly in their notational symbols.

What are Some of the Advantages and Disadvantages of the RD Design?

When we talk about advantages or disadvantages of a research design, it is important to ask "Compared to what?" Most of the time, the RD design is compared with other designs which are similar to it, usually other pre-post multi-group designs. Obviously, all pre-post designs have certain advantages and disadvantages. For instance, they all enable one to estimate "gain" or "change" over time. On the other hand, they tend to be more expensive and time-consuming than one-shot studies like surveys. The RD design shares all of the advantages and disadvantages that are common to all before-and-after two-group designs. But to really understand the advantages

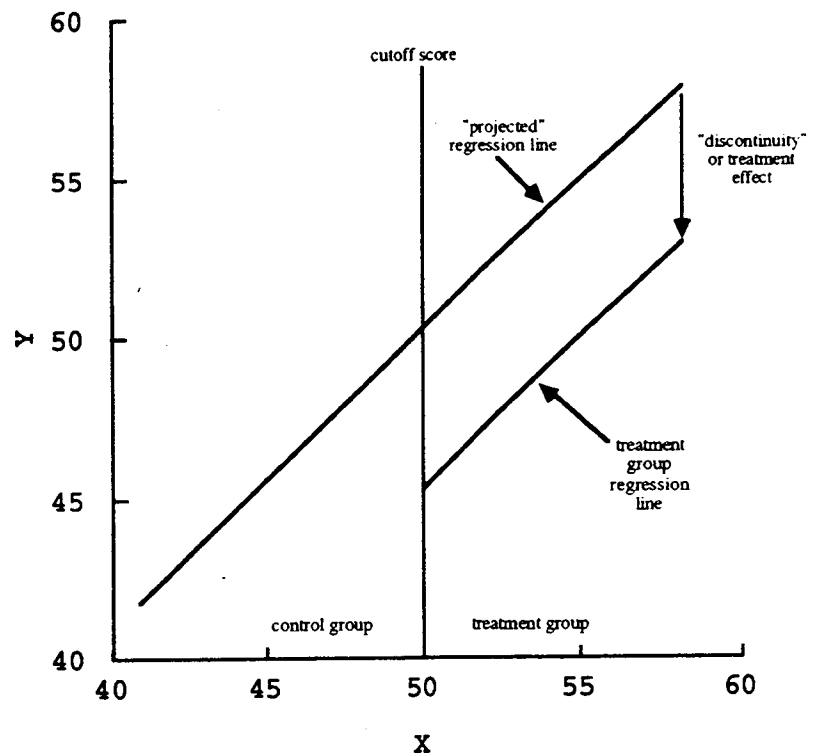


Figure 8. Graph of only the regression lines for an RD design.

and disadvantages of the RD design, it is important to grasp how it differs from its most similar competitors.

There are two types of pre-post two-group designs which are usually compared with the RD design. All three have the same design structure:

$$\begin{array}{ccc} O & X & O \\ O & & O \end{array}$$

and all of them can be distinguished from each other by the method used to assign persons to treatments or programs. In a *randomized experiment*, sometimes also called the *randomized clinical trial*, persons are assigned to treatments randomly or by chance, as in a lottery. As we have seen, in the RD design, persons are assigned solely on the basis of a cutoff rule. Finally, in a *nonequivalent group design*, the assignment of persons is not controlled by the researcher. For instance, if patients are free to choose which group they would like to be in, we have a nonequivalent group design. Comparisons of these three designs involve trade-offs -- there is no simple way to say that one design of the three is always going to be preferable to the others. However, research methodologists do have some general preferences which hold true.

Most researchers will prefer to use a randomized experiment over any other design. A randomized experiment assures that both groups will be comparable prior to the treatment. Any differences which are observed after the treatment are more likely to be attributable to the treatment rather than to any initial differences between groups. But randomized experiments are not perfect designs for all situations. They require that persons in the study be randomly assigned. Thus, the clinicians and administrators have to give up control of treatment assignment during the course of the study. Furthermore, clients and/or clinicians may not like the treatment which is randomly allocated. An extremely ill patient may not be satisfied with being assigned to a no-treatment control group. Similarly, a mildly ill person may not be willing to be assigned to an untested, high-intensity treatment or program. In some cases, clients will drop out of the study, seek other treatments, or clinicians will provide compensating treatments to help them out. Any of these conditions could jeopardize the original random assignment and lead to incorrect conclusions in the randomized experiment.

The nonequivalent group design is generally thought to be the weakest of the three. Because the researcher does not control the assignment of the treatment, there may be an infinity of factors which differ between the groups prior to the treatment. After treatment is given, it is likely to be extremely difficult to tell whether any differences on outcome measures are due to the treatment or to other differences which existed between the groups. Nevertheless, the nonequivalent group design is probably the most widely used of the three because it occurs naturally. If we want to test the effect of a new treatment, we can give the treatment to all of the patients we have access to and then try to find another "comparable" group which didn't get the new treatment. For instance, we might go to a nearby agency that works with similar clients but which is not trying out the new treatment. Although it is the easiest design to implement, one can seldom be sure that the two groups are comparable or equivalent except with regards to the treatment -- the term "nonequivalent" group design serves to remind us of this built in frailty.

The RD design is generally viewed as somewhere between these two designs in its ability to estimate the effects of a treatment. It is generally preferred to the nonequivalent group design because, even though the groups are not equivalent in pretest ability (after all, they're made very nonequivalent on the pretest by the cutoff rule!), the researcher controls this nonequivalence exactly and is able to adjust for it in the analysis of the data. In fact, if you think back to the discussion of regression lines as shown in Figures 6 through 8, you can see that the RD design is based on a notion of equivalence -- it assumes that the two groups would be well described by the same regression line if the treatment is not given (as in Figure 6). If we observe a discontinuity

between the regression lines (as in Figure 7 or 8) this is likely to be due to the treatment rather than to any other factor.

The RD design is generally considered weaker than the randomized experiment, and the latter is generally preferred. For one thing, we know that we need more patients in an RD study than in a randomized experiment in order to attain the same level of statistical precision. But there are situations in which a randomized experiment may not be feasible. Most often, the randomized experiment is threatened for ethical reasons -- it requires an arbitrary assignment to treatments which may not meet the wishes or needs of the client or clinician. Clinicians don't randomly assign treatments (at least we hope not!), they try to assign the treatment based on the client's clinical symptoms and expressed need. Radical surgery would not usually be prescribed for minor illnesses and mild therapies would not usually be given for life-threatening illness. In a randomized experiment, one may have to assign people contrary to desirable clinical practice in order to test out the new treatment.

The real power of the RD design lies in its closer correspondence with accepted clinical practice. In the RD design, one can give the treatment to those who need it most or those who are most willing to take the risk of trying it. Similarly, with the RD design, you don't have to assign potentially costly treatment to those who are less sick or less willing to risk it.

Thus, the RD design is a compromise design. It has some of the scientific quality of a randomized experiment in that the researcher controls the assignment to treatments and is more likely to be able to make a valid assessment of the effectiveness of the program. But it also has some of the desirable qualities of good clinical practice, allowing the researcher to assign new treatments to those who might need them or be willing to risk them. As with any compromise, the RD design will never be a perfectly satisfying choice. But it is a realistic alternative when one is unable to randomly assign, and has more scientific credibility than nonequivalent self-selecting assignment procedures.

Summary

The Regression-Discontinuity (RD) design is a pre-post two-group design which can be used to assess the effects or outcomes of treatments or programs. The RD design is characterized by its unique method of assigning persons to receive the treatment or program -- persons are assigned solely on the basis of a cutoff score on a pretreatment measure. That is, all persons scoring on one side of the pre-selected cutoff score are assigned to receive the new treatment or program while those scoring on the other side are assigned to a no-treatment or traditional treatment control group. The RD design minimally requires that one have a pretreatment score, a group membership score (1 if treated; 0 if control) and a posttreatment score. The RD design is so named because when the treatment is effective, there is a discontinuity in the regression lines at the cutoff point. Generally, RD designs are not considered to be as strong as randomized experiments, but are considered superior to nonequivalent group designs. However, in situations where randomized experiments are not feasible -- this will most often be for ethical reasons -- the RD design should be considered the strongest alternative method available.

Bibliography

Good introductory discussions of the RD design can be found in:

Trochim, W. (1984). *Research Design for Program Evaluation: The Regression-Discontinuity Approach*. Beverly Hills, CA: Sage Publications.

Judd, C.M., and Kenny, D.A. *Estimating the effects of social interventions*. Cambridge, MA: Cambridge University Press, 1981.

Glossary

assignment criterion	The method or standard used to assign persons to treatments or programs in a study. There are three major types of assignment: random, rule-based and uncontrolled. Assignment by a cutoff score, as in the RD design, is one of the most common rule-based assignment criteria.
before-and-after two-group design	A research design where all participants are divided into two groups (usually treatment and control) by an assignment criterion and are measured before and after the administration of the treatment.
bivariate plot	An x-y plot showing the relationship between two variables. Each variable is represented using one of the two axes. A person's score on each variable is represented by a single point placed at the intersection of that person's score on each variable. In pretest-posttest designs, the pretest is usually graphed on the x-axis and the posttest on the y-axis.
control group	A group that does not receive the treatment or program being investigated. Sometimes the control group receives no treatment, sometimes it gets a pseudo (placebo) treatment, and sometimes it gets the standard treatment (in a relative comparison study).
cutoff score	An assignment criterion which characterizes the RD design. The cutoff score forms the rule for assigning persons to treatments. Those who score on one side of the cutoff on a pretreatment measure are assigned to the treatment, while those scoring on the other side are assigned to the control group.
design notation	A shorthand method for describing a research design. Although design notations differ slightly from one text to another, most indicate time from left to right, treatment groups on separate lines, measures or observations with an 'O,' treatments with an 'X' and the assignment criterion with a special symbol (e.g., 'R' for random assignment, 'C' for cutoff assignment, 'N' for nonequivalent assignment).
line of best fit	Another name for a regression line. This terminology indicates that it is desirable that the regression analysis yield a line which describes the relationships among variables well.
negative correlation	A negative relationship between two variables. In a negative correlation, high values on one variable are associated with low values on the other and vice versa.

nonequivalent group design	A before-and-after two-group design where the assignment criterion is not controlled by the researcher. In a nonequivalent group design, for instance, persons may self-select which treatment they would like to participate in.
null case	The case where no treatment is administered or where the treatment, if administered, has no effect.
positive correlation	A positive relationship between two variables. In a positive correlation, high values on one variable are associated with high values on the other and low values on one variable are associated with low values on the other.
posttest	A measure or observation which is collected after the treatment or program has been administered. Usually a posttest is measured because it is thought that it will show the effects of a program or treatment.
pretest	A measure or observation which is collected before the treatment or program has been administered.
program group	A group that receives the treatment or program being investigated.
randomized clinical trial	The term typically used for a randomized experiment in medical research.
randomized experiment	A before-and-after two-group design where persons are randomly assigned to either the treatment or control group.
regression analysis	A statistical method used to describe relationships between variables. The results of a regression analysis can be displayed as a regression line which summarizes a bivariate relationship.
regression line	A line which can be drawn after conducting a regression analysis. The regression line is ideally a good summary description of the relationships among variables.
regression-discontinuity design	A before-and-after two-group design where persons are assigned to either the treatment or control group solely on the basis of a cutoff score assignment criterion on a pretest measure.
research design	The description of how the measures, participants, and groups are related to each other over the course of a research study.
treatment effect	The demonstrated effect of the treatment or program.

treatment group

A group that receives the treatment or program being investigated.