

CHAPTER III

RESEARCH METHODOLOGY

A. Research Design

This research used an experimental design. The experimental design are procedures in quantitative research in which the investigator determines wheather an activity or materials make a difference results for participants (Creswell,2012). The purpose of the experiment is to assign participants to experience it, have some participants experience something different, and then establish whether the participants who got used to the concept, practice, or process outperformed the non-participants on some outcome. To determine a potential cause-and-effect relationship between independent and dependent variables, researchers conduct experiments.

The type of experimental design that used in this research is Quasi - experimental research. Quasi-experimental design that used is pre- and posttest (Creswell,2012). The experimental and control treatments were given to intact groups by the researcher, who then gave a pretest to both groups, conducted experimental treatment activities with the experimental group exclusively, and then gave a posttest to evaluate the distinctions between the two groups. For more details, see the following table.

Table 3.1
Quasi Pretest–Posttest Design

Group	Pretest	Independent	Posttest
Experiment	<i>Y 1</i>	<i>X</i>	<i>Y 2</i>
Control	<i>Y 1</i>	-	<i>Y 2</i>

This design was to evaluate the effectiveness of Mentimeter for tenth grade of senior high school. The researcher was given a pretest for both classes, then gave treatment by using Mentimeter to the experimental group only, and then gave a posttest to both groups. If the experimental group shows significantly greater achievement on the posttest, the researcher can conclude that the Mentimeter is effective.

The objective of the research was to investigate the significant effect of Mentimeter on students' reading comprehension ability. In this research, there are two variables: the students' reading comprehension is a dependent variable, while Mentimeter is an independent variable.

B. Variable of the Research

In this study, there are two variables. They are Independent Variable (*x*) and Dependent Variable (*y*). Variable is a characteristic or attribute of an individual or an organization that researchers can measure or observe and varies among individuals or organizations studied (Creswell, 2012).

1. The independent variable

An independent variable is an attribute or characteristic that influences or affects dependent variable (Creswell, 2012). One independent variable must be the treatment variable. The independent variable of this research was the use of Mentimeter in teaching reading comprehension.

2. The dependent variable

A dependent variable is an attribute or characteristic that is dependent on or influenced by the independent variable (Creswell, 2012). The dependent variable of this study was the students' achievement in the reading comprehension after using Mentimeter.

C. Population of the Research

Population is a group of individuals who have the same characteristic (Creswell 2012). The population of this research is the tenth grade of MAN 2 Pesisir Selatan as population it has 9 classes. Population data in this study can be seen through the following table.

Table 3. 2
Population of The Research

No	Class	Students
1.	XA	34
2.	XB	34
3.	XC	36
4.	XD	34
5.	XE	35
6.	XF	35
7.	XG	35
8.	GH	35

9.	XI	34
Total	9 Classes	314 Students

D. Sample & Sampling Technique

The sample represents an aspect of the population's number and characteristics. Cluster random sampling is used by researchers when the population has already been sorted into natural, preexisting groupings. A cluster could be a district, school, classroom, metropolitan statistical area, city zone area, neighborhood block, or street, among other things. A researcher may select a particular street and randomly sample those who live there (Schreiber 2011). The researcher chose two classes from nine by entering the names of the nine classes and selecting numbers and obtaining the first name class as the experiment class (XA), which has 34 students, and the second name as the control class (XB), which also has 34 students.

Table 3. 3
Sample of The Research

No	Class	Students
1.	XA	34
2.	XB	34
Total	2 Classes	68 Students

E. Instrument of the Research

Instruments are very important in the research to help researcher to collect data and information. The researcher used a test as instrument in this research. The type of test that used in this research was multiple choices. Huges

(1998) adds the explanation above that possible solution for measure of reading academic texts are multiple choice question, short answer, summary, close information transfer.

The researcher had tried out 40 multiple choice questions in different classes, namely in X.C. After carrying out the try out result from the participant, the researcher tested the validity of the questions based on the result of try out. The researcher used 25 multiple choice after validity and reability tested. According to Brown (2019) Indicators of reading comprehension is main idea, expressions/idioms/phrases in context, inference (implied detail), grammatical features, detail (scanning for a specifically stated detail), excluding facts not written (unstated details), supporting idea, vocabulary in context. For more details, the instrument can be seen in the appendix.

Table 3. 4
The Blue Print of Reading Comprehension Test

Variable	Indicators	Number
Reading Comprehension	Students can find the main idea	3, 6, 13,16, 23
	Students can find inference (implied detail)	1, 2, 7, 10, 17, 18, 19,21, 24, 25
	Students can identifying grammatical features	4,
	Students can find the excluding facts not written (unstated details)	5, 8, 9, 12, 15, 22
	Students can find the vocabulary in context	14,20

F. Data Collecting Technique

To collect data for this study (pre-test and post-test). The test use to determine the students' level of reading comprehension. The test is administered twice, the first is a pre-test, and the second was a post-test to students' reading comprehension at the tenth grade of MAN 2 Pesisir Selatan.

1. Pre-test

The pre-test is a test that will be conducting before to some treatments. The students answered multiple choices questions in the pre test based on the indicator of reading comprehension by the syllabus in the school. Pre-test implemented in November 19 2024 with 25 multiple choices. Students have 60 minutes to answer the questions.

2. Post-test

The post-test is a test that was conducting after some treatments. The students was answering 25 multiple choice questions in the post test based on the indicator of reading comprehension in Desember, 4 2024. The students have 60 minutes to answer to the questions.

G. Technique of Data Analysis

In analyzing the data, the researcher will apply some steps as follows:

1. Descriptive analysis

Descriptive analysis includes sample numbers, lowest and maximum scores, as well as mean and standard deviation. A descriptive analysis was conducted using students' pre-test and post-test results. In this

example, the researcher employed two approaches, one manually and one using the SPSS version, which can be used to produce descriptive data, and the SPSS program was used to calculate all of the data acquired in this study, namely the paired sample T-test. The researchers aim to discover how utilizing Mentimeter affects reading comprehension. Researchers began by administering a pretest. Following the pretest, the researcher administered the treatments three times. Researchers use Mentimeter in their process text. Students receive a post-test. Finally, the results of the two tests were compared to determine whether there whether there was a significant difference between the pre-test and post-test.

- a. To know the interval data of the pretest and posttest the Researcher calculate the value of the tank by using the following formula:

$$I = R / K$$

Where:

I = Interval

R = Range ($X_{\max} - X_{\min} + 1$)

K = Number of Class ($1 + (3.3 * \log N)$)

- b. To find out the improvement of percentage

$$\% = \frac{X_2 - X_1}{X_1} \times 100\%$$

Where:

% : The percentage of improvement

X₂ : The total score of Post-Test

X1 : The total score of Pre-Test

- c. To find out the mean score of the students test, the researcher used the formula:

$$\bar{x} = \frac{\sum FiXi}{\sum fi}$$

Where:

X = Mean Score

$\sum xi$ = Total score

$\sum fi$ = The number of students

- d. Determining of standart deviation score of variable

$$SD = \sqrt{\frac{\sum x^2 - \frac{(\sum X)^2}{N}}{N-1}}$$

Where:

X = Mean Score

$\sum x$ = The sum of the difference score

N = The number of students

2. Normality Test

The normality test examined whether the pre-test and post-test data were normally distributed. When the p-value is greater than 0.05, the data can be considered normal. The Kolmogrov-Smirnov test with SPSS 26 was used to verify the normality of reading test scores in the form of pre-test and post-test data.

a. Pre-Test Standard Normality

In this research, the pre-test of the experimental and control classes was conducted on November 19. The aim of this test is to know the students' reading comprehension ability in procedural text before treatments were given. To quantify their reading comprehension ability in procedural text, the students were given 25 multiple-choice questions.

A normality test can be performed using either the Kolmogorov-Smirnov or Shapiro-Wilk methods. The distinctions between these normalcy tests take the shape of research samples. Kolmogorov-Smirnov is employed when there are more than 50 samples in a study sample. Shapiro-Wilk is employed when the sample size is 50 or fewer. To determine data normality, a significant value > 0.05 indicates normal data, whereas a significant value < 0.05 indicates non-normal data. The data in this study was processed using SPSS's Kolmogorov-Smirnov test. The results are as follows:

Table 3.5
Test of Normality

One-Sample Kolmogorov-Smirnov Test		Unstandardized Residual
N		34
Normal	Mean	.0000000
Parameters ^{a,b}	Std. Deviation	11.50283869
	Absolute	.077

Most Extreme	Positive	.077
Differences	Negative	-.070
Test Statistic		.077
Asymp. Sig. (2-tailed)		.200 ^{c,d}

- a. Test distribution is Normal.
- b. Calculated from data.
- c. Lilliefors Significance Correction.
- d. This is a lower bound of the true significance.

Based on the result of the normality test, the significance value is 0.200, which is $0.200 > 0.05$; it can be concluded that the data were normal.

a. Post-test Standar Normality

To determine the normality of the data, the researcher utilized Kolmogrov-Smirnov test to evaluate the sample data. A significant value more than 0.05 indicates normal data, while a significant value less than 0.05 indicates non-normal data. The result of the normality test can be seen from the table below.

Table 3.6
Test of Normality

One-Sample Kolmogorov-Smirnov Test

		Unstandardized Residual
N		34
Normal	Mean	.0000000
Parameter	Std. Deviation	14.22854607
s ^{a,b}		
Most	Absolute	.123
Extreme	Positive	.123

Differences	Negative	-0.075
Test Statistic		.123
Asymp. Sig. (2-tailed)		.200 ^{c,d}

- a. Test distribution is Normal.
- b. Calculated from data.
- c. Lilliefors Significance Correction.
- d. This is a lower bound of the true significance.

Based on the table above, the K-S score of the sample was 0.200. It showed $0.200 > 0.05$, which means that the data was normal. It can be concluded that the data from the post-test was normal.

3. Homogeneity Test

The homogeneity test aims to determine whether the variance is homogeneous or not in the two data groups. This test can be carried out using the Levene test using SPSS 26 software, so that a significance value is obtained. Data is categorized as normal if the significance is greater than 0.05

a. Pre-Test Standard Homogeneity

Following the normality test, the homogeneity test was used to determine whether the two populations share the same variance. The homogeneity test of the two variances between the experimental and control classes was calculated using the Levene statistic test. The results are as follows:

Table 3.7
Test of Homogeneity
Test of Homogeneity of Variances

	Levene Statistic	df1	df2	Sig.
Based on Mean	3.085	1	66	.084
Based on Median	2.659	1	66	.108
Based on Median and with adjusted df	2.659	1	60.627	.108
Based on trimmed mean	3.151	1	66	.081

The Levene's statistic for equality of variances above reveals that the data's significance score was 0.084, which is greater than 0.05. Because the data exceeded 0.05, we can assume that the data were homogeneous or equal.

b. Post-Test Standard Homogeneity

The homogeneity test determines if the samples come from the same variance or not. Homogeneous data is defined to contain the same variance. The researcher used SPSS version 26 to analyze data homogeneity through the One-Way ANOVA test. If $\alpha \text{ sig} > 0.05$, the two samples are considered homogeneous. If $\alpha \text{ sig} < 0.05$, the two samples are considered inhomogeneous. The results are as follows:

Table 3.8
Test of Homogeneity

Test of Homogeneity of Variances		Levene Statistic	df1	df2	Sig.
result readin g	Based on Mean	3.667	1	66	.060
	Based on Median	3.576	1	66	.063
	Based on Median and with adjusted df	3.576	1	61.906	.063
	Based on trimmed mean	3.622	1	66	.061

Due to the Levene's statistic for equality of variances, the data had a significance score of 0.60. Therefore, it indicates 0.060 is greater than 0.05. Because the data exceeded 0.05, it can be concluded that the data was homogeneous.

4. T- test

The T-test was used to analyze the data. The-test was divided into two types: independent t-test and paired sample t-test. In learning process, the T-test was used to determine whether the different data from pre-test and post-test. The researcher used paired sample t-test. Paired sample t-test was a test that used two paired samples. The t-test was concluding through SPSS 26. In this research, the researcher chose the t-test with the paired sample t-Test formula, namely :

$$t = \frac{\bar{D}}{\frac{SD}{\sqrt{N}}}$$

Note :

$\bar{D}$ = Standard Mean Deviation

SD = Standard Deviation of $\bar{D}$

N = Number of Sample

Before the research used paired sample T-test, there are the steps that the research did, as follow :

- a. The research determined the mean of pretest (x) and mean of post-test (y)

The formula :

$$x = \frac{\sum x}{N}$$

$$y = \frac{\sum y}{N}$$

UIN IMAM BONJOL
PADANG

Notes :

$\sum x$ = Total scores of pre-test

$\sum y$ = Total score of post-test

N = Total number of sample

- b. After that, the research calculated the mean of differentiate pre-test and post-test with the formula :

$$Md = \frac{\sum d}{N}$$

Note :

Md = The mean of differential pre-test and post test

$\sum d$ = The sum of different between pre-test and post test

N = Total number of students

- c. Next, the researcher counted the standard deviation, the researcher used the formulation

$$s = \sqrt{\frac{\sum x^2 - \frac{(\sum x)^2}{N}}{N-1}} \qquad s = \sqrt{\frac{\sum y^2 - \frac{(\sum y)^2}{N}}{N-1}}$$

Notes :

S = Standard Deviation

$\sum x^2$ = Sum of pre-test quadrate score

$\sum x$ = Sum of pretest score

$\sum y^2$ = Sum of Pre-test quadrate score

$\sum y$ = Sum of Post –test score

N = Number of Students

- d. The researcher calculated the total number of quadrate deviation with the formula used :

$$\sum X^2 d = \sum d^2 - \frac{(\sum d)^2}{N}$$

Notes :

$\sum x^2 d$ = Total of quadrate deviation

$\sum d$ = Sum of different between pre-test and post test

N = Number of student

- e. The researcher figure out the t-test with formula :

$$t_{count} = \frac{Md}{\sqrt{\frac{\sum x^2 d}{N(N-1)}}}$$

Notes :

Md = Mean different of pre-test and post-test

$\sum x^2 d$ = Total of quadrate deviation

N = total of students

- f. Last the researcher determined the t- table distribution with significant 5%

$$1. \quad df = N - 1$$

Notes :

df = Degree of freedom

N = Number Students (Gay, 1981: 366)

When T-test result Positive (+)

1. If the ρ -value is higher than sig $\alpha = 0.05$ (5%) then the alternative hypothesis is accepted
2. If the result indicates ρ -value or sig (2-talled) lower than the significance level of sig $\alpha = 0.05$ (5%), then the null hypothesis is accepted.

When T-test result negative (-)

1. If the result indicates ρ -value or sig (2-talled) lower than the significance level of sig $\alpha = 0.05$ (5%), then the alternative hypothesis is accepted

2. If the result indicates p -value or sig (2-tailed) higher than the significance

5. Validity and Reliability

a. Validity

The validity test indicates the level of validity of a test. High validity is achieved when the measurement function is accurate and findings align with the intended purpose. To determine if each question is legitimate or invalid, if the r count $>$ r table with a significance level of 0.05, the instrument is considered valid. If R count $<$ R table with a significance level of 0.05, the instrument is considered invalid.

b. Reliability

This test assesses the consistency of questionnaire results when used repeatedly. In other words, it can be utilized repeatedly by the same or other researchers and produce the same results. A Cronbach's alpha value of 0.6 or higher indicates reliability. For measuring equipment to demonstrate high reliability, the reliability coefficient (α) must approach one. A measuring instrument is deemed reliable if its alpha coefficient (α) is larger than 0.6 and less reliable if it is less than 0.2.