

BAB 2

TINJAUAN PUSTAKA

2.1 Profil Puskesmas Watubelah

Penelitian ini dilaksanakan di Puskesmas Watubelah. Berikut merupakan uraian profil dari Puskesmas Watubelah.

2.1.1 Sejarah Puskesmas Watubelah

Puskesmas Watubelah merupakan Unit Pelaksana Teknis Daerah yang merupakan bagian dari Dinas Kesehatan Kabupaten Cirebon yang bertanggung jawab penuh terhadap seluruh upaya pelaksanaan kesehatan yang berada pada wilayah kerja di Kecamatan Sumber Kabupaten Cirebon. Berdasarkan Peraturan Bupati Cirebon Nomor 43 Tahun 2018 tentang Organisasi, Fungsi, Tugas Pokok, dan Tata Kerja Unit Pelaksana Teknis Daerah Pada Dinas Kesehatan Kabupaten Cirebon, Puskesmas Watubelah termasuk dalam kategori Puskesmas Non Rawat Inap dan memiliki karakteristik wilayah yang termasuk kedalam Puskesmas Kawasan Perkotaan. Selain itu, berdasarkan Keputusan Bupati Cirebon Nomor 440/Kep.366/Dinkes/2019 tentang Penetapan Pusat Kesehatan Masyarakat Mampu Pelayanan Obstetri Neonatal Emergensi Dasar (Puskesmas Mampu PONED), Puskesmas Watubelah telah dinyatakan menjadi bagian dari Puskesmas mampu PONED.

2.1.2 Visi, Misi dan Tujuan Puskesmas Watubelah

Secara geografis, Puskesmas Watubelah terletak pada daerah perkotaan di Kelurahan Watubelah, Kecamatan Sumber, Kabupaten Cirebon, dengan letak astronomis pada 6°44'1"S, 108°29'48"E. Puskesmas Watubelah sendiri memiliki luasan wilayah sebesar 94,7 Km². Secara administratif Puskesmas Watubelah berbatasan langsung dengan wilayah kerja kecamatan dan kabupaten lain dengan batas-batas wilayahnya adalah sebagai berikut :

Sebelah Utara : Kecamatan Weru

Sebelah Timur: Kecamatan Tengah Tani

Sebelah Selatan: Kecamatan Sumber

Sebelah Barat : Kecamatan Plumbon

Puskesmas Watubelah terletak di Jalan Tangkil Gede nomor 05 Kelurahan Watubelah, Kecamatan Sumber Kabupaten Cirebon. Secara administratif Puskesmas Watubelah memiliki 5 kelurahan sebagai batasan wilayah kerja, dimana terdapat 43 RW dan 199 RT. Dimana kelurahan Watubelah menjadi kelurahan dengan jarak tempuh terdekat, yaitu sejauh 0,5 km, dimana Kelurahan Watubelah juga menjadi lokasi keberadaan Puskesmas Watubelah. Sedangkan kelurahan terjauh memiliki jarak tempuh sejauh 3 km, yang terletak di kelurahan Kenanga.

1. Visi

Menjadi Puskesmas yang bermutu dalam memberikan pelayanan untuk mewujudkan masyarakat yang sehat dan mandiri di wilayah kerja Puskesmas Watubelah.

2. Misi

1. Memberikan pelayanan kesehatan dasar sesuai kebutuhan masyarakat;
2. Meningkatkan kualitas sumber daya manusia sesuai kompetensinya untuk meningkatkan mutu pelayanan;
3. Meningkatkan kerjasama lintas sektoral dan partisipasi masyarakat dalam pembangunan kesehatan.

3. Tujuan

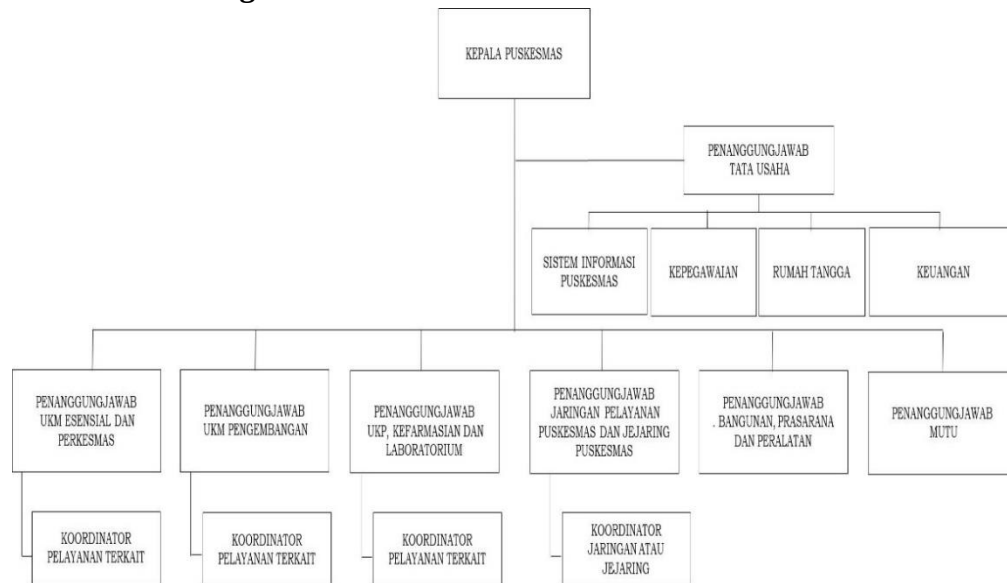
Meningkatkan kemandirian masyarakat sesuai dengan visi misi Puskesmas Watubelah

2.1.3 Logo Puskesmas Watubelah



Gambar 2. 1 Logo Puskesmas Watubelah

2.1.4 Struktur Organisasi



Gambar 2. 2 Struktur Organisasi

1. Jobdesk

Berikut merupakan jobdesk dari Penanggung jawab Kefarmasian selaku pihak yang akan dibantu dalam penelitian ini:

Penanggung jawab Kefarmasian bertanggung jawab kepada pasien yang berkaitan dengan Sediaan Farmasi dengan maksud mencapai hasil yang pasti untuk meningkatkan mutu kehidupan pasien. Selain itu, Pelayanan Kefarmasian adalah tolok ukur yang dipergunakan sebagai pedoman bagi tenaga kefarmasian dalam menyelenggarakan pelayanan kefarmasian. Penanggung jawab Kefarmasian juga memiliki tanggung jawab dalam persediaan Farmasi seperti obat, bahan obat, obat tradisional dan kosmetika di Puskesmas Watubelah.

2.2 Landasan Teori

Landasan teori ini berisikan teori-teori pendukung yang digunakan dalam proses analisis dan implementasi pada permasalahan yang diangkat dalam penelitian ini.

2.2.1 Data

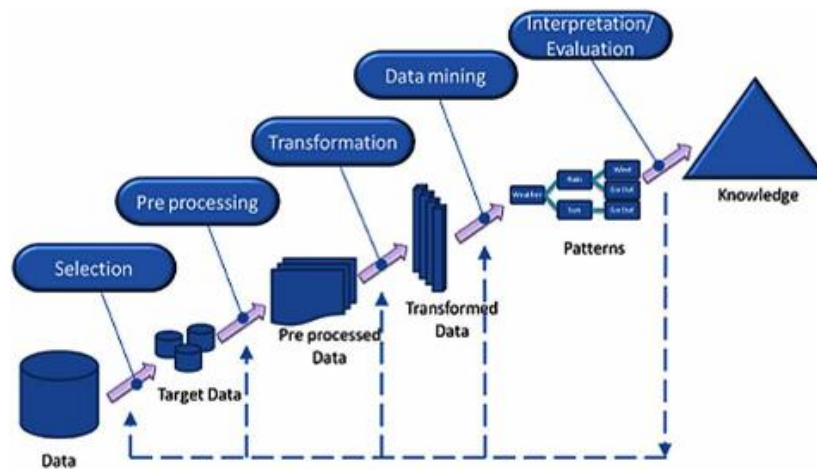
Data merupakan suatu informasi yang memberukan gambaran yang lebih luas terkait suatu kondisi dan keadaan yang biasanya terdiri dari fakta-fakta. Pada proses pengambilan keputusan maupun kebijakan umumnya seseorang akan menggunakan data sebagai bahan pertimbangan untuk menganalisis, menggambarkan maupun

menjelaskan suatu keadaan[5]. Data biasanya bersumber dari proses pencarian dan pengamatan dari sumber-sumber tertentu. Data masih bersifat fakta mentah yang selanjutnya perlu dilakukan beberapa proses agar data tersebut dapat memberikan informasi yang lebih mudah dipahami.

Berdasarkan sifatnya data terbagi menjadi dua, yaitu data kualitatif dan data kuantitatif. Data kualitatif seringkali ditunjukkan dengan data yang tidak berbentuk angka. Adapun data kuantitatif merupakan suatu data yang biasanya dinyatakan dalam bentuk angka. Sedangkan menurut sumbernya data haruslah mengacu pada sumber perolehan data, yaitu data eksternal dan data internal. Data internal biasanya bersumber dari suatu perkumpulan yang berbetuk kelompok dan organisasi. Sedangkan data eksternal adalah suatu data yang sumbernya berasal dari luar kelompok atau organisasi tertentu. Menurut cara perolehannya, data dibedakan menjadi data primer dan data sekunder. Data primer merupakan data yang diperoleh dengan cara dikumpulkan dan diolah langsung dari objeknya oleh peneliti. Sedangkan data sekunder adalah data yang didapatkan dalam bentuk data jadi yang sebelumnya telah diolah oleh pihak terkait. Sedangkan data menurut waktu pengumpulannya, dibedakan menjadi data *cross section* dan data berkala atau data *time series*. Data *cross section* ialah data yang diperoleh pada waktu dan periode tertentu. Sedangkan data berkala ialah data yang didapatkan dan dikumpulkan dari waktu ke waktu.

2.2.2 Knowledge Discovery in Databases

Knowledge Discovery in Databases (KDD) merupakan proses untuk memperoleh pengetahuan yang ada pada data, yang kemudian akan digunakan untuk menekankan kepada hasil yang memiliki kualitas dengan tingkatan yang tinggi dengan metode tertentu dengan menggunakan Data Mining. Dengan menggunakan suatu ukuran dan batas pada suatu basis data, pengetahuan akan diekstrak, yang kemudian akan dilanjutkan dengan proses preprocessing, sub sampling dan transformasi dari basis data tersebut[6]. Gambar 2.3 menunjukkan tahapan dari proses KDD.



Gambar 2.3 Knowledge Discovery Data

Berikut adalah penjelasan dari masing-masing tahapan :

1. *Selection*

Pada proses *selection* akan dilakukan pembuatan dan pengumpulan data set yang kemudian akan diolah untuk mendapatkan pengetahuan atau *knowledge discovery*.

2. *Pre processing*

Tahapan *Pre processing* akan melakukan proses pembersihan data dengan menghilangkan duplikasi data, melakukan pengecekan data yang dinilai tidak konsisten dan memperbaiki data yang dinilai memiliki kesalahan.

3. *Transformation*

Pada tahapan ini akan terjadi perubahan pada data yang telah dipilih, dimana data tersebut dipastikan telah sesuai untuk selanjutnya dilakukan proses *data mining*. Jenis dan pola informasi menjadi dua hal yang akan dipertimbangkan pada tahapan ini yang selanjutnya akan dicari dalam *database*.

4. *Data Mining*

Pada tahapan *data mining* akan dicari pola ataupun informasi yang dinilai menarik pada data yang telah terpilih dengan menggunakan teknik ataupun metode tertentu. Adapun teknik, metode ataupun algoritma yang ada pada *data mining* sangat bervariasi. Dengan keberagaman metode dan algoritma tersebut,

sangat penting untuk memilih metode dan algoritma yang tepat yang sesuai dengan tujuan dan proses KDD secara menyeluruh.

5. *Interpretation / Evaluation*

Tahapan *interpretation* merupakan bagian dari KDD yang terdiri atas pemeriksaan apakah pola maupun informasi yang didapatkan bertentangan dengan fakta maupun hipotesa yang telah dibuat sebelumnya. Selanjutnya, pola informasi yang didapatkan dari proses *data mining* akan ditambahkan dalam bentuk yang mudah dipahami oleh pihak yang berkepentingan.

Proses KDD bertujuan untuk mencari informasi-informasi yang berharga, pola yang ada di dalam data, yang melibatkan algoritma untuk mengidentifikasi pola pada data. meringkas proses KDD dari berbagai step, yaitu: seleksi data, pra-proses data, transformasi data, data mining, dan yang terakhir interpretasi dan evaluasi. diawali dari data mentahan, dan terus dilakukan hingga berakhir pada informasi atau pengetahuan yang sudah diolah[7].

2.2.3 *Clustering*

Clustering adalah suatu teknik atau metode yang banyak digunakan untuk mengelompokkan data yang ada menjadi beberapa kelompok atau kluster yang dibutuhkan berdasarkan kesamaan antar data. Dengan adanya *clustering* kita dapat menemukan struktur data yang sebelumnya tidak diketahui[8]. *Clustering* merupakan kategori unsupervised learning dengan menggunakan data yang tidak mempunyai label dan kemudian akan ditemukan struktur, pola dan hubungannya. [9]. Apabila data yang ada pada suatu *cluster* memiliki kemiripan yang tinggi dan memiliki ketidakmiripan yang tinggi pula dengan data yang terdapat pada *cluster* yang lain akan menghasilkan *Clustering* dengan kualitas yang baik. Hal ini menunjukkan bahwa hasil *cluster* akan bergantung pada algoritma yang akan digunakan serta karakteristik dari data itu sendiri.

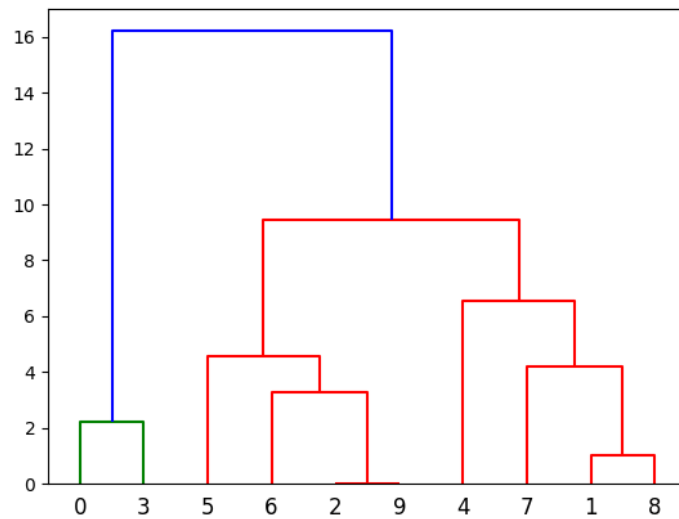
2.2.4 *Agglomerative Hierarchical Clustering*

Agglomerative adalah suatu pendekatan pengelompokkan dari hierarki yang diawali dengan melakukan penggabungan objek yang ada pada satu *cluster* yang

terpisah yang kemudian akan membentuk *cluster* yang semakin besar[10]. *Agglomerative Hierarchical Clustering* adalah adalah suatu teknik dari *hierarchical clustering* yang kemudian akan menghasilkan satu *cluster* tunggal dengan menggabungkan N buah *cluster*. Pada tahapan awal dari metode ini akan diletakkan setiap data sebagai satu *cluster* tersendiri atau *atomic cluster* yang kemudian akan digabungkan dengan beberapa *cluster* tersebut dan kemudian akan menjadi sebuah *cluster* tunggal[11].

Adapun langkah-langkah dari proses algoritma *Agglomerative Hierarchical Clustering* :

1. Melakukan perhitungan matrik jarak antar data dengan menggunakan rumus *Euclidean Distance*.
2. Menggabungkan kelompok yang memiliki karakteristik terdekat menjadi satu dengan menggunakan perhitungan *ward*.
3. Menghitung kembali jarak antar *cluster* yang baru terbentuk dengan titik data lainnya.
4. Mengulangi langkah 2 sampai dengan langkah 3 hingga terbentuk satu *cluster* untuk semua data. Dimana langkah ini akan diulang hingga mampu membentuk satu *cluster* yang sesuai dengan ide dasar dari *Agglomerative Hierarchical Clustering*.
5. Melakukan evaluasi *cluster* terbaik dengan menggunakan perhitungan *Davies Boulding Index* atau DBI.
6. Membentuk dendogram untuk memvisualisasikan *cluster* yang terbentuk. Gambar 2.4 merupakan contoh bentuk dendogram.



Gambar 2.4 Dendogram

Visualisasi dari proses *clustering* tidak selalu menggunakan dendogram. Banyak jenis-jenis visualisasi yang dapat menampilkan hasil dari proses penemuan pengetahuan menggunakan *data mining*, salah satunya adalah *Scatter Plot*. Contoh bentuk *scatter plot* dapat dilihat pada **Error! Reference source not found..**

2.2.5 Perhitungan Jarak

Metode *Ward* merupakan sebuah Teknik yang ditemukan pada tahun 1963 oleh Ward. Dimana metode *ward* merupakan suatu Teknik yang mampu menghasilkan *cluster* dengan mengurangi variansi yang ada pada *cluster* tersebut. Dengan menggunakan metode *ward*, setiap jarak Euclid kuadrat ke rata-rata

cluster harus dihitung. Dimana metode ini tidak akan menghitung jarak antar kelompok, melainkan akan membentuk kelompok-kelompok dengan memaksimalkan kehomogenan dari kelompok tersebut[12]. Setiap *cluster* rata-rata dari seluruh variable yang ada akan dihitung, lalu akan dilakukan perhitungan untuk menganalisis *cluster* dengan menggunakan rumus berikut:

$$Ward(S_i, S_j) = \frac{N_{S_i} N_{S_j}}{N_{S_i} + N_{S_j}} d(c_{S_i}, c_{S_j}) \quad (2.1)$$

Nilai N_{S_i} dan c_{S_i} masing-masing mewakili kardinalitas dan *centroid* dari *cluster* S_i , sementara N_{S_j} dan c_{S_j} mewakili kardinalitas dan *centroid* dari *cluster* S_j , dan $d(c_{S_i}, c_{S_j})$ adalah fungsi yang mengembalikan jarak antara *centroid* masing-masing *cluster* S_i dan S_j .

2.2.6 Perhitungan Jarak

Jarak merupakan suatu hal yang akan dipertimbangkan dalam proses analisis data. Dengan menggunakan jarak, kita dapat mengukur kedekatan antara dua objek ataupun titik data. Dimana dalam menentukan titik data yang terdekat dapat menggunakan kueri tertentu. Oleh karena itu, diperlukan perhitungan antara titik kueri dan titik lainnya. Dengan menggunakan matrik jarak kita dapat membentuk batas keputusan. Dimana untuk mengetahui titik serupa yang terdekat, dapat dilakukan menggunakan perhitungan jarak, seperti *Euclidean Distance*, *Hamming Distance*, *Manhattan Distance* dan *Minkonskidistance*[13].

1. *Euclidian Distance*

Euclidean distance merupakan salah satu metode perhitungan jarak yang memiliki 2 buah titik yang berada pada *Euclidean Space*. Dimana *euclidean space* diperkenalkan pada tahun 300 B.C.E. (*before the Common Era*) oleh seorang matematikawan yang berasal dari Yunani, Euclid. Dimana Euclid membuat metode ini untuk mempelajari hubungan antara sudut dan jarak. Metode *Euclidean* ini akan berkaitan dengan Teorema Phytagoras yang biasany akan diterapkan pada 1,2 dan 3 dimensi. Metode ini juga dinilai lebih sederhana apabila diterapkan pada dimensi yang lebih tinggi[14].

Berikut adalah rumus perhitungan *Euclidean Distance* [15] :

$$dist(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (2.3)$$

Dengan D adalah jarak antara titik pada data x dan titik data y, dimana $x = x_1, x_2, \dots, x_i$ dan $y = y_1, y_2, \dots, y_i$ dan j mempresentasikan nilai atribut serta p merupakan dimensi atribut.

Keterangan:

d = jarak antara x dan y

k = variabel data (setiap data)

n = banyak parameter (jumlah data)

x = data testing

y = data training

2. Davies Bouldin Index

Metode Davies Bouldin Index temukan oleh David L. Davies dan Donald W. Boulding pada tahun 1979. Dimana Davies dan Boulding menemukan metode untuk mengevaluasi hasil dari pengelompokkan yang telah didapatkan untuk mengetahui jumlah *cluster* yang paling optimal. Metode ini dinilai mampu mengukur validitas serta menghasilkan jumlah *cluster* yang paling optimal. Dimana pada proses pengelompokkan, kohesi diartikan sebagai jumlah dari kedekatan data t ke *cluster* lain. Dalam DBI, akan ditemui Sum of Square within Cluster (SSW) yang merupakan sebuah persamaan untuk mencari nilai dari matriks kohesi pada *cluster* ke-i. Nilai SSW ini dapat dihitung dengan rumus:[16] :

$$SSWi = \frac{1}{m_i} \sum_{j=i}^{m_i} d(x_j c_i) \quad (2.4)$$

Selain itu, ada persamaan Sum of Square Between Cluster (SSB) yang digunakan untuk mengetahui nilai separasi antar cluster. Nilai SSB dapat dihitung dengan rumus[16] :

$$SSBi,j = d(C_i, C_j) \quad (2.5)$$

Setelah memperoleh nilai kohesi dan separasi menggunakan persamaan sebelumnya, selanjutnya dilakukan pengukuran rasio untuk mengetahui nilai

perbandingan antara cluster ke-i dan cluster ke-j. Cluster akan dikatakan baik jika memiliki nilai kohesi sekecil mungkin dan nilai separasi yang sebesar mungkin. Nilai rasio dihitung dengan rumus [16] :

$$R_{i,j} = \frac{SSW_i + SSW_j}{SSB_{i,j}} \quad (2.6)$$

Nilai rasio yang telah diperoleh akan digunakan untuk mencari nilai Davies Bouldin Index. Nilai Davies-Bouldin Index dapat dihitung dengan rumus :

$$SSWi = \frac{1}{K} \sum_{l=i}^K \max_{l \neq j} (R_{i,j}) \quad (2.7)$$

Nilai yang didapat dari rumus tersebut merupakan hasil yang menjadi ukuran validitas dari cluster yang diuji. Semakin kecil nilai yang didapatkan (non-negatif ≥ 0), maka semakin baik cluster yang diperoleh.

2.2.7 Unified Modelling Language (UML)

Unified Modeling Language (UML) merupakan sebuah metodologi untuk mengembangkan sistem yang berorientasi pada objek dan juga termasuk alat untuk membangun pengembangan sistem. Dengan UML kita dapat mendokumentasikan, menspesifikasikan serta membangun perangkat lunak yang merupakan bahasa spesifikasi standar. UML juga dapat memvisualisasikan, menspesifikasikan, membangun serta mendokumentasikan sebuah bahasa yang berbentuk grafik ataupun gambar menjadi sebuah sistem pengembangan *software* berbasis *Object Oriented* atau OO. UML sendiri juga akan memberikan standar penulisan sebuah sistem blue print, yang meliputi konsep bisnis proses, penulisan kelas-kelas dalam bahasa program yang spesifik, skema database, dan komponen-komponen yang diperlukan dalam sistem software[17] .

1. Use Case

Use case diagram adalah sebuah pemodelan yang dibuat untuk melakukan sistem informasi yang kemudian akan kita rancang. Dengan menggunakan *use case* kita dapat mendeskripsikan hubungan atau korelasi antara satu peran maupun lebih dengan sistem informasi yang kita rancang. Selain itu kita juga dapat menggunakan *use case* untuk mengetahui fungsi apa saja yang terdapat pada sistem berita dan siapa saja yang berhak untuk menggunakan fungsi-fungsi itu[18].

2. Activity Diagram

Activity diagram merupakan suatu diagram yang menggambarkan laur kerja ataupun aktivitas dari sebuah sistem atau proses bisnis yang akan dilakukan. *Activity diagram* juga akan menjelaskan aktivitas sistem yang dilakukan oleh sistem bukan aktor. Selain itu, diagram ini juga akan menampilkan tindakan dan bagian dasar dari proses tansisi yang kemudian akan memicu tindakan penyelesaian yang berasal dari sumber. *Activity diagram* kerap kali dihubungkan dengan *flowchart* yang menggambarkan proses yang terjadi antara aktor dan sistem[19].

3. Class Diagram

Class diagram merupakan suatu diagram yang berisi tentang penjelasan mengelai kelompok objek yang ada dengan properti, operasi (prilaku), serta relasi yang sama. Dimana diagram kelas adalah salah satu jenis diagram yang dapay menggambarkan susunan kelas dan hubungan yang ada antar kelas tersebut yang kerap kali digunakan dalam pemodelan perangkat lunak. Diagram kelas ini merupakan bagian *Unified Modeling Language*, dimana UML ini sendiri merupakan bagian dari bahasa standar untuk mendokumentasikan dan memodelkan perangkat lunak[20].

4. Sequence Diagram

Sequence diagram adalah suatu diagram yang mampu mendeskripsikan interaksi antar objek dan mampu mengidentifikasikan komunikasi yang ada diantara objek tersebut. *Sequence diagram* mampu menunjukkan interaksi proses satu sama lain dengan fungsi menunjukkan proses satu sama lain dan kelas yang saling terlibat dalam perencanaan diantara benda-benda yang diperlukan[21].

2.2.8 Python

Python adalah bahasa pemrograman serbaguna yang mudah dipelajari baik untuk pemula maupun ahli. Kita dapat menggunakan Python karena memiliki perpustakaan serbaguna yang siap digunakan untuk semua jenis aplikasi, mulai dari pemrograman statistik hingga pembelajaran mendalam. Keuntungan Python adalah bahasa pemrograman yang ideal dengan sintaksis yang mudah dipahami (Junaidi,

2023). Python menyediakan berbagai paket visualisasi. Hal ini menjadikan Python penting tidak hanya untuk analisis data tetapi juga untuk semua ilmu data, karena ia memiliki antarmuka visual yang memungkinkan Anda mengakses dan menggunakan data dengan lebih mudah dengan membuat berbagai bagan, grafik, dan visual interaktif. Oleh karena itu, Python memungkinkan pengguna untuk lebih mudah memahami data[22].

2.2.9 Website

Website adalah situs informasi yang disediakan di Internet dan dapat diakses oleh siapa saja di seluruh dunia. *Website* diartikan sebagai kumpulan halaman yang memuat informasi data digital berupa teks, gambar, animasi, suara, video atau kombinasinya, yang didistribusikan melalui koneksi Internet dan tersedia untuk diakses dan dilihat oleh siapa saja dunia. Halaman situs web dibuat dalam bahasa standar HTML. Skrip HTML ini diterjemahkan oleh *browser web* dan ditampilkan sebagai informasi yang dapat dibaca siapa pun.[23].

2.2.10 Streamlit

Streamlit adalah kerangka kerja Python sumber terbuka. Streamlit juga menyediakan kemampuan untuk mengubah skrip menjadi aplikasi web yang dapat dilihat dalam hitungan menit. Hal ini membuatnya sangat mudah untuk membuat aplikasi berbasis web, terutama bagi orang-orang yang memiliki sedikit pengetahuan tentang desain pengembangan *web front-end*.

[24]

2.2.11 Black Box Testing

Pengujian perangkat lunak dalam pengertian spesifikasi fungsional, bukan pengujian desain atau kode program, untuk menentukan apakah fungsionalitas, masukan, dan keluaran perangkat lunak memenuhi spesifikasi yang diperlukan., Metode pengujian black box mudah digunakan karena hanya memerlukan batas bawah dan batas atas data yang diharapkan. Perkiraan kumpulan data pengujian dapat dihitung berdasarkan jumlah kolom data masukan yang akan diuji, aturan masukan yang harus dipenuhi, dan kasus batas atas dan bawah yang dipenuhi[25].