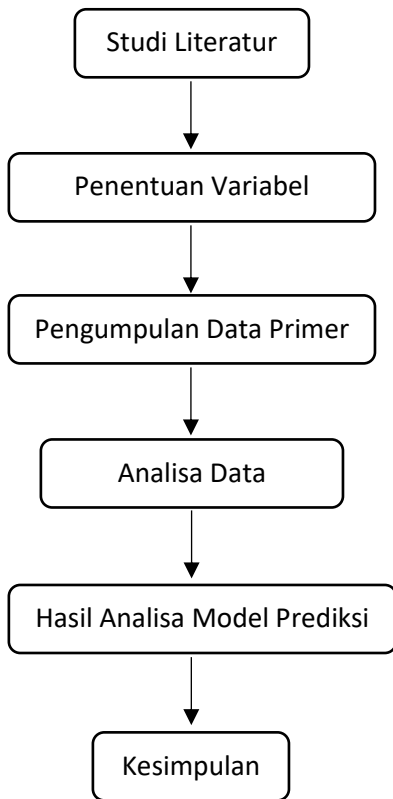


3 METODOLOGI PENELITIAN

3.1 Kerangka Penelitian

Pada bab ini akan dibahas mengenai metodologi penelitian yang akan dilakukan untuk mencapai tujuan penelitian yang telah ditetapkan. Ada beberapa tahapan penelitian yang dilakukan oleh peneliti tertera pada Gambar 3.1.



Gambar 3.1. Kerangka Penelitian

Penelitian diawali dengan melakukan studi literatur, hal ini dilakukan untuk mempelajari lebih lanjut mengenai penelitian yang akan dilakukan serta merangkum permasalahan dan melihat hasil penelitian yang pernah dilakukan sebelumnya. Studi literatur juga dilakukan untuk menentukan beberapa faktor atau variabel yang pernah digunakan pada penelitian sebelumnya dan kemudian akan digunakan menjadi variabel dalam penelitian ini. Setelah terpilih beberapa faktor, dilakukan pengumpulan data primer. Data primer yang dikumpulkan berisi variabel – variabel dari proyek yang pernah dikerjakan. Kemudian, dilanjutkan dengan pengolahan data disini proses pengolahan data dilakukan dengan menyusun data agar mempermudah dalam pelaksanaan pembentukan model yang akurat. Setelah data terbentuk maka akan mulai dilakukan pembentukan model yang dibantu oleh program IBM SPSS. Dari hasil analisa akan

mendapatkan pembentukan model prediksi. Dari model yang didapatkan akan diambil kesimpulan yang dapat menjawab rumusan masalah dari penelitian.

3.2 Jenis Penelitian

Penelitian ini merupakan jenis penelitian kuantitatif. Penelitian kuantitatif adalah penelitian dengan uji hipotesis yang dikembangkan berdasarkan teori, kemudian akan diuji berdasarkan data empiris yang telah dikumpulkan. Jenis penelitian kuantitatif diharapkan dapat dipakai untuk menjawab tujuan penelitian yaitu membentuk model dan mengevaluasi metode yang paling baik untuk memprediksi estimasi biaya konstruksi pada proyek konstruksi.

3.3 Populasi dan Sampel

3.3.1 Populasi

Menurut Sugiyono (2017), populasi adalah wilayah generalisasi yang terdiri atas objek atau subjek yang mempunyai kuantitas dan karakteristik tertentu yang ditetapkan oleh peneliti untuk dipelajari dan kemudian ditarik kesimpulan. Populasi yang akan digunakan adalah proyek konstruksi bangunan yang ada di pulau Jawa.

3.3.2 Sampel

Sampel merupakan bagian dari populasi yang dianggap dapat mewakili populasi yang diteliti. Sampel yang digunakan akan diambil menggunakan metode *non-probability sampling* dengan *purposive sampling*. *Non-probability sampling* adalah metode pengambilan sampel dengan tidak memberi peluang yang sama bagi setiap anggota populasi untuk dipilih kembali menjadi sampel. *Purposive sampling* merupakan metode pengambilan sampel didasarkan seleksi khusus dengan kriteria yang telah ditentukan. Pengambilan kriteria ditentukan terlebih dahulu, untuk menemukan data informasi yang sesuai dengan hasil yang diinginkan. Kriteria sampel adalah data proyek konstruksi yang telah terselesaikan. Dalam penelitian ini target sampel yaitu 100 proyek pembangunan yang telah dilakukan di area pulau Jawa.

3.4 Sumber dan Jenis Data

Jenis data yang digunakan dalam penelitian ini berupa data primer. Menurut Umar (2014), data primer merupakan data yang didapat dari sumber pertama baik dari individu perseorangan maupun objek populasi berupa hasil pengisian kuesioner atau formulir yang dilakukan oleh peneliti. Pada penelitian ini formulir yang diberikan kepada responden akan berisi

variabel penelitian yang telah ditetapkan berdasar studi literatur. Formulir akan diberikan kepada responden baik secara langsung maupun secara online via email maupun Google Forms.

3.5 Metode dan Prosedur Pengambilan Data

Dalam penelitian ini metode pengambilan data dilakukan dengan memberikan form kuisioner lewat google form maupun diisi secara langsung. Kuisioner akan dibagikan kepada responden yaitu kepada pihak kontraktor. Data yang terkumpul akan diolah dan diklasifikasikan sesuai dengan kebutuhan. Adapun penyajian dalam kuisioner akan berisi beberapa data teknis dari hasil pembangunan proyek konstruksi yang pernah diselesaikan sebelumnya.

3.6 Definisi Operasional Variabel

Pengumpulan data dilakukan dengan melakukan pencarian data pembangunan proyek konstruksi yang telah selesai dibangun pada beberapa kontraktor. Data akan diklasifikasikan sesuai dengan faktor-faktor atau variabel seperti pada Tabel 3.1.

Tabel 3.1.

Variabel yang Digunakan Dalam Penelitian

No.	Variabel	Satuan / Klasifikasi
x1	Luas proyek / bangunan	m ²
x2	Skala kontraktor	1 = Kecil
		2 = Menengah
		3 = Besar
x3	Jumlah tingkat bangunan	unit
x4	Jumlah kolom struktur	unit
x5	Luas lansekap	m ²
x6	Total tinggi bangunan	m
x7	Volume beton	m ³
x8	Volume profil baja	kg
x9	Tahun pembangunan	2000,2001, . . .
x10	Klasifikasi bangunan	1 = Rumah Tinggal
		2 = Komersial
		3 = Gudang
		4 = Pabrik

Tabel 3.1.

Variabel yang Digunakan Dalam Penelitian (lanjutan)

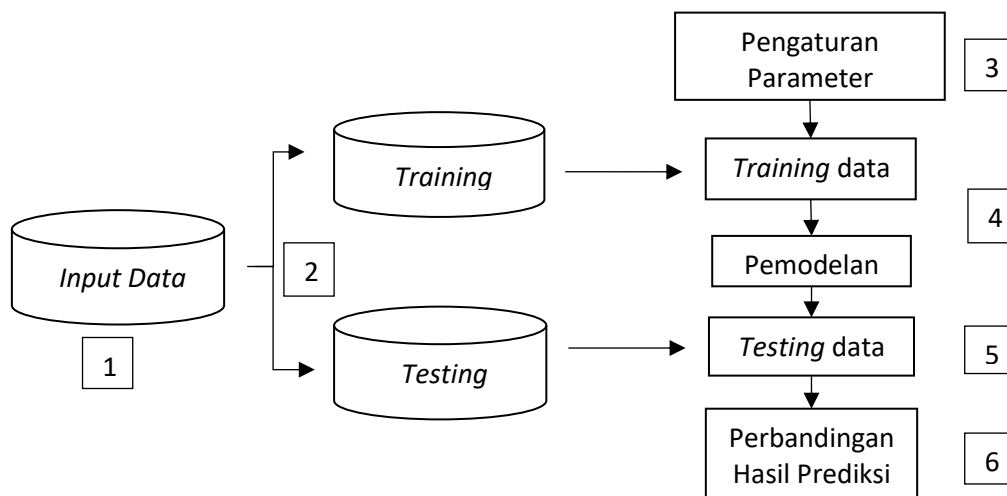
No.	Variabel	Satuan / Klasifikasi
x11	Tipe pondasi	1 = Pondasi Setempat
		2 = Pondasi Strouss
		3 = Tiang Pancang
x12	Tipe pemilik	1 = Swasta
		2 = Pemerintah
x13	Tipe Kontrak	1 = Lump Sum
		2 = Unit Price
x14	Lokasi proyek	1 = Kota
		2 = Daerah
x15	Metode Pemilihan Kontraktor	1 = Tender
		2 = Pemilihan
		3 = Penunjukkan
y	Biaya awal proyek (Target)	Rp. (rupiah)

3.7 Analisa dan Pengolahan Data

Pada sub-bab ini dijelaskan mengenai cara pengambilan, pengolahan dan penyajian hasil prediksi dari data yang didapatkan. Pengolahan data dan pengaplikasian metode prediksi akan dilakukan menggunakan IBM SPSS 18.0. Gambar 3.2 menunjukkan gambaran singkat bagaimana proses prediksi ini dilakukan beserta urutan langkah-langkah penelitian:

Adapun langkah prosesnya adalah sebagai berikut:

1. Langkah pertama yang dilakukan yaitu tahapan pengumpulan data, pada tahap ini akan dilakukan pengambilan data dari beberapa kontraktor dan proyek konstruksi yang telah berlangsung sesuai variabel yang telah ditentukan dari hasil pembelajaran atau studi literatur
2. Setelah terkumpul data maka akan dilakukan pemisahan data menjadi *training* dan *testing* secara acak. Pembagian jumlah sampel yang pas juga menjadi salah satu faktor penting dalam menghasilkan model prediksi yang tepat. Menurut Hoang & Nguyen (2018), dan Hoang & Tran (2019) komposisi *training* dan *testing* yang baik adalah 70:30.
3. Selanjutnya, dilakukan pengaturan parameter agar mendapat model prediksi yang akurat.
4. Langkah keempat, dilakukan proses *training* yang nantinya akan menghasilkan permodelan untuk masing-masing metode.
5. Langkah kelima, dilakukan proses *testing* untuk mengetahui apakah pemodelan yang dihasilkan dapat dengan akurat memprediksi biaya proyek pembangunan sesuai aktual.



Gambar 3.2. Diagram Alir Proses Prediksi

6. Langkah terakhir yaitu, 4 metode tersebut akan dibandingkan mana metode prediksi yang paling akurat hasil prediksinya dengan cara melihat hasil dari 5 metode evaluasi.

3.7.1 Pengaturan Parameter

Untuk mendapatkan hasil dari prediksi yang tetap sesuai dan tidak abstrak maka perlu diberikan pengaturan terhadap parameter. Selain itu, adanya pengaturan parameter bertujuan untuk mengevaluasi pengaturan parameter metode prediksi yang dapat menghasilkan prediksi yang paling akurat. Pengaturan parameter diaplikasikan untuk 3 jenis AI yang digunakan pada penelitian ini, yaitu: *artificial neural network* (ANN), *support vector machine* (SVM), dan *classification and regression tree* (CART). Pada metode LR tidak diperlukan adanya pengaturan parameter tambahan, hal ini karena LR merupakan regresi linier. Pengaturan parameter yang digunakan metode LR hanya pengaturan parameter standar dari *software*.

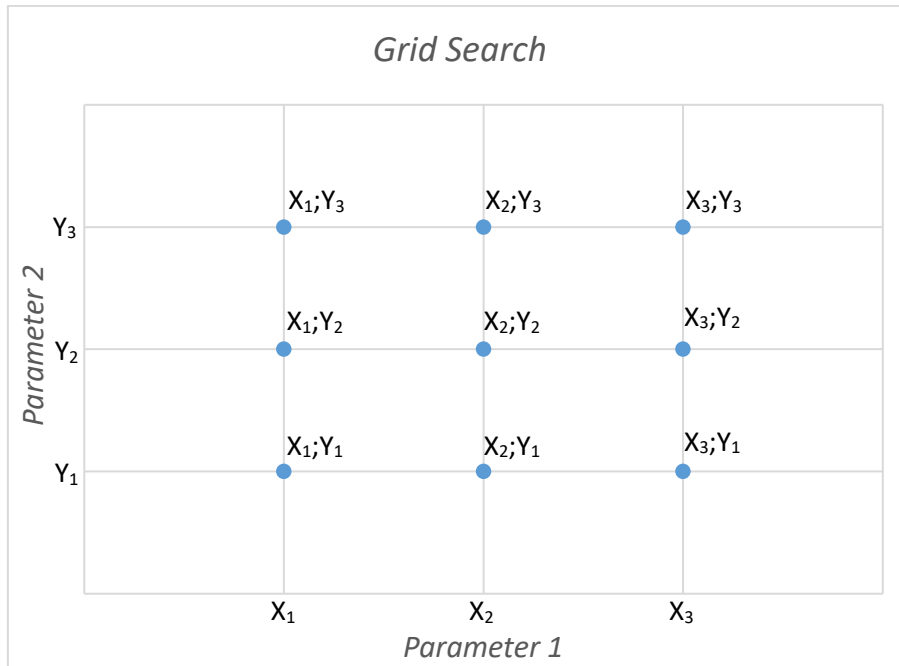
Pengaturan parameter yang digunakan pada penelitian ini adalah standar dari IBM SPSS Modeler 18.0 dan terdapat beberapa perubahan pengaturan parameter yang menganut kaidah *grid search*. Pengaturan parameter masing- masing metode terlampir pada Tabel 3.2. Nomor 1 pada akhiran nama metode prediksi menjelaskan bahwa metode tersebut berjalan sesuai dengan pengaturan parameter standar.

Tabel 3.2.

Pengaturan Parameter Standar pada IBM SPSS Modeler 18.0

Metode Prediksi	Parameter	Pengaturan
ANN	<i>ANN Model</i>	<i>Multilayer Preceptron</i>
	<i>Maximum Training Time</i>	<i>15 Minutes</i>
	<i>Maximum Number of Cycle</i>	1000
	<i>Hidden Layer</i>	{1, 2}
	<i>Number of Hidden Neuron</i>	<i>Automatic</i>
SVM	<i>Epsilon</i>	1
	<i>Kernel</i>	RBF
	<i>Regularization Parameter, C</i>	<i>Automatic</i>
	<i>Kernel parameter, γ</i>	<i>Automatic</i>
CART	<i>Levels Below Root</i>	5
	<i>Max Surrogates</i>	5
	<i>Change in Impurity</i>	0.0001
	<i>Impurity Measure</i>	Gini
	<i>Prune Tree</i>	TRUE
LR	<i>Collinearity Diagnostics</i>	FALSE
	<i>Regression Coefficients</i>	TRUE
	<i>Singularity Tolerance</i>	0.0001

Grid search adalah metode sederhana yang dapat digunakan untuk optimasi suatu parameter. Metode ini menggabungkan 2 parameter dengan bantuan produk kartesius untuk menemukan kombinasi parameter yang baru. Pada ANN, parameter yang diubah adalah *number of neuron hidden layer 1* dan *number of neuron hidden layer 2*. Sedangkan untuk metode SVM, parameter yang diubah adalah nilai dari C (*regularization parameter*) dan nilai dari γ (*kernel parameter*). Gambaran mengenai *grid search* dapat terlihat pada Gambar 3.3 berikut ini.



Gambar 3.3. Gambaran Umum *Grid Search*

Untuk metode CART, parameter yang diubah adalah parameter “*levels below root*”. Pengaturan “*levels below root*” standar IBM SPSS adalah 5 root, yang akan diubah-ubah untuk pemodelan CART lainnya. Kombinasi parameter pada CART dapat dilihat pada Tabel 3.3.

Tabel 3.3.

Pengaturan Parameter CART

Nama	Parameter
CART2	<i>levels below root</i> = n_1
CART3	<i>levels below root</i> = n_2
CART4	<i>levels below root</i> = n_3

Semua pengaturan parameter ini diaplikasikan pada data yang ada, setelah itu dilakukan evaluasi model prediksi dengan parameter manakah yang punya performa paling bagus dalam memprediksi biaya konstruksi.

3.7.2 Proses *Training*

Proses *training* adalah proses dimana sebagian data diolah untuk melakukan pembelajaran tersendiri, nantinya data tersebut akan digunakan sebagai model prediksi. Sebagian besar pendekatan dengan proses *training* cenderung bisa memprediksi data yang berlaku tidak sesuai atau terletak jauh dari garis regresi pada umumnya. Proses penyesuaian

dengan data nyata ini dinamakan *fit the data*. Banyaknya data yang digunakan untuk proses ini sebanyak 70% dari total data. Hasil dari proses training ini adalah model prediksi setiap metode yang merepresentasikan setiap *dataset*.

3.7.3 Proses *Testing*

Proses *testing* adalah proses implementasi model yang sudah didapatkan dalam training pada 30% sisa data yang belum pernah dilakukan pada proses *training*. Jika pola data yang dimiliki dalam proses *testing* sesuai dengan data *training*, maka nilai prediksi yang dihasilkan akan menyerupai dengan data *training*-nya.

3.7.4 Perbandingan Hasil Prediksi

Setelah ditemukan hasil *testing* dari setiap metode prediksi yang ada, nilai-nilai korelasi, dan *error* yang sudah ada dibandingkan dalam satu tabel dengan setiap metode yang digunakan. Pada penelitian ini batas *error* yang akan digunakan yaitu pada MAPE 50%.