

Chapter 1

Introduction

1.1 Context

Global Navigation Satellite Systems (GNSS), based on trilateration, are used for positioning and navigation of the vehicle. They are accessible sources all over the world used for finding the position of the user by estimating the distance between the receiver and the satellites. At least 4 satellites are needed for a PVT solution. Once the position of each satellite is known, we could get an estimation of the user position. However, there still exist errors as the signal propagates such as ionospheric and tropospheric errors, clock errors, multipath, interferences, etc. In harsh environments (the presence of large buildings for instance), these systems are not efficient. These systems can be ineffective especially in urban environments because the signals sent by satellite are reflected, diffracted, and refracted leading to an inaccurate position estimation, adding a large bias to the actual position of the user (which is not desirable). Also, the signals from the satellites are lost, for example, when driving through a tunnel, and then the receiver cannot estimate its position and therefore, track the vehicle trajectory. For simple calculations, we can use least squares estimation techniques which are not very accurate. In reality, when we implement a PVT solution for the estimation of vehicle/user trajectory, our problem is non-linear. Thus, it becomes important to consider these characteristics to correctly model the system accurately.

1.2 Objective

The idea of this project is to implement different estimators. The project uses 3 estimation techniques namely, Least Square Estimation (LSE), Extended Kalman Filter (EKF), and Particle Filter (PF), along with comparison of their performance. A detailed study of PF is performed by studying its various types. Next, the concept of Map-Matching is implemented, where the estimated position of the user obtained after the implementation of the estimator (PF) is merged with the map information. These maps can be used to provide a priori information on the receiver trajectory and act as assistance data. The project implements map matching with PF as it can handle non-standard priors and likelihood. Two methods of map matching are proposed: the first method uses the distance between the map and particles to update the observation equation and the second method draws a proposal distribution from intersecting points of a circle with the center at the estimated position, at the previous epoch. The project aims to develop a positioning algorithm that merges the observations of a GNSS receiver and a map based on a particle filter in order to improve the positioning of a vehicle on the road, especially in an urban environment. It is developed assuming a 2D space (x, y position of the satellite). Later, this would be used to develop MAGNITUDE, the TELECOM/SIGNAV library for GNSS data processing.

Chapter 2

State of Art about GNSS

Global Navigation Satellite System (GNSS) refers to a constellation of satellites that transmit positioning and timing data from space to GNSS receivers via signals. The concept of GNSS positioning is the idea of computation of the user position, velocity, and time using this available data. It uses the measurements of time-of-arrival from multiple satellite signals to find estimates of distances to the satellites. For this, first, the satellite position is computed, before the clock correction. To that extent, several parameters (such as semimajor axis, corrected mean motion, delta t) are extracted from the data [1]. Then, rotation arguments and corrected orbital parameters are computed in order to obtain the satellite position in the Earth-Centred Earth Fixed (ECEF) frame at the time of transmission. To have satellite ECEF coordinates in reception time, the Sagnac effect is taken into account. Finally, the satellite positions in ECEF are converted in the Latitude-Longitude-Altitude coordinate system (WGS-84). This estimated satellite position is an input to the estimator, as well as the estimated satellite clock bias. The output is a rough estimate of the user position. This rough estimate allows computing tropospheric and ionospheric errors. Finally, we use this estimator to obtain a fine and correct estimate of the user position by taking into account all possible errors. In the section for estimation techniques, we find the use of an estimator. We can implement simple estimators like the Least Square Estimator (LSE), the Extended Kalman Filter, or the Particle Filter. As discussed, LSE techniques have very low accuracy. This accuracy degrades even more in an urban environment, due to multipath. To solve this problem and to ensure that we can have accurate positioning, Bayesian filters like Kalman filters(KF), Extended Kalman filters(EKF), Particle filters(PF), etc may be used. This project aims to compare various estimators and therefore, draws a performance comparison. Methods for improving the position estimation after the implementation of estimators (PF) will also be discussed.

In the rest of this work, we will assume a 2D positioning problem where all the various terms, except the receiver clock bias and noise, are perfectly corrected. In this context, the equation of the pseudorange of the i -th satellite is expressed as [1]:

$$\tilde{\rho}^i = \sqrt{(x - x_{sat})^2 + (y - y_{sat})^2} + b_u + n^i, \quad (2.1)$$

where

- (x, y) denotes the position of user in the 2D frame
- (x^i, y^i) denotes the position of the i -th satellite expressed in the same 2D frame
- b_u is the satellite clock bias
- n^i represents the measurement noise.

Chapter 3

Problem Formulation

In this chapter, we will present the formulation of our problem. This includes the different estimation techniques used for positioning. The concept and formulae of some useful estimation techniques will be discussed that were implemented on MATLAB. This chapter also introduces the concept of map-matching and the approaches taken in view of this project.

3.1 Estimation Techniques

In this section, we will present the estimators and their formulations. Classical estimators like least square estimators will be presented along with Bayesian estimators like Kalman filters, Extended Kalman Filters, and Particle filters.

3.1.1 Least Square Estimation

Ordinary Least Squares

The first estimation technique for positioning is the least square estimation (LSE) or simply the ordinary least squares. By definition, the LSE minimizes the squared norm of estimation residuals, i.e., the difference between the real observations and their predictions for a given value of the state vector. In the case of a linear model we thus have:

$$\min_{\mathbf{x}} \|\mathbf{y} - \mathbf{H}\mathbf{x}\|_2^2 \quad (3.1)$$

leading

$$\hat{\mathbf{x}}_{LS} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \mathbf{y}. \quad (3.2)$$

This method involves building a measurement matrix from the partial derivatives of the geometric distances between the satellite and the user with respect to each coordinate axis. The difference between the previous value of the state vector and its new value is obtained by the formation of its pseudoinverse matrix and factoring it into the measurement equation with all the corrections of the pseudorange. Therefore, their new value needs to be added to the previous one as an iteration process, which is interrupted when the said difference is low enough or when the number of iterations for the same epoch is high enough. Since the geometric range and the user clock bias depends on the current estimation of the user position, the newly obtained values are reported back to the initial phase of the process if it has not been stopped by the aforementioned conditions, and the process is then repeated.

Mathematically, let us consider the state vector to be estimated as $\mathbf{x}$, which gathers the position and the clock bias. Then the measurement vector is modeled as $\mathbf{y} = \mathbf{h}(\mathbf{x})$, where $\mathbf{y}$ represents

the GNSS pseudoranges or measurements. The noise vector is assumed Gaussian with zero mean and with covariance matrix $\mathbf{R}$, i.e., $\mathbf{n} \sim \mathcal{N}(0, \mathbf{R})$. The resulting measurement equation is: [9]

$$\mathbf{y} = \mathbf{h}(\mathbf{x}) + \mathbf{n} \quad (3.3)$$

But for the GNSS problem, the above measurement equation is not linear, therefore we can not use ordinary least squares and therefore, need to resort to the non-linear least squares (explained below).

Non-linear Least Squares

When the measurement model is not linear, ordinary least squares can not be used instead we use non-linear least squares, where the positioning solution is based on the linearization of the measurement model around a supposed state vector, say, $\hat{\mathbf{x}}_0$:

$$\mathbf{h}(\mathbf{x}) = \mathbf{h}(\hat{\mathbf{x}}_0) + \frac{\partial \mathbf{h}}{\partial \mathbf{x}} \mathbf{h}(\hat{\mathbf{x}}_0) \cdot (\mathbf{x} - \mathbf{x}_0), \quad (3.4)$$

where $\frac{\partial \mathbf{h}}{\partial \mathbf{x}} \mathbf{h}(\hat{\mathbf{x}}_0) = \mathbf{H}$, called as the Jacobian matrix of $\mathbf{h}(\mathbf{x})$. Thus the measurement model can be re-written as:

$$\mathbf{y} = \mathbf{h}(\hat{\mathbf{x}}_0) + \mathbf{H} \cdot (\mathbf{x} - \mathbf{x}_0) + \mathbf{n} \quad (3.5)$$

OR

$$\Delta \mathbf{y} = \mathbf{H} \Delta \mathbf{x} + \mathbf{n} \quad (3.6)$$

Hence, this reformulated model is linear between:

- (a) The transformed measurement, $\Delta \mathbf{y} = \mathbf{y} - \mathbf{h}(\hat{\mathbf{x}}_0)$, and
- (b) The transformed state, $\Delta \mathbf{x} = \mathbf{x} - \hat{\mathbf{x}}_0$

Therefore, the Least Square Estimator formula can be then re-used as:

$$\Delta \hat{\mathbf{x}}_{LS} = (\mathbf{H}^T \mathbf{H})^{-1} \mathbf{H}^T \Delta \mathbf{y} \quad (3.7)$$

This implies that the final estimate is:

$$\hat{\mathbf{x}} = \hat{\mathbf{x}}_0 + \Delta \hat{\mathbf{x}}_{LS} \quad (3.8)$$

Usually, this method is applied iteratively in order to refine the linearization point, $\hat{\mathbf{x}}_0$, and this iterative least square method has been implemented in the code script as well.

3.1.2 Kalman Filter (KF)

In Kalman Filter, we are in a Bayesian framework, and hence we are looking for the posterior distribution, $p(\mathbf{x}_k | \mathbf{y}_{1:k})$ and the state vector. Each epoch is represented by ' k '. The following is the state space model with first order Markov assumption for the state model, which is true for KF and also for EKF and PF (described in later sub-sections):

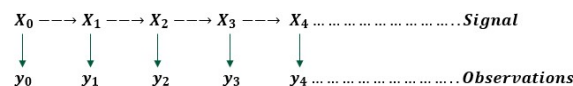


Figure 3.1: System model (adapted from [5])

In this study, the velocity and clock drift is added to the state that makes our state vector as:

$\mathbf{x} = [\mathbf{x}_k \quad \dot{\mathbf{x}}_k \quad \mathbf{y}_k \quad \dot{\mathbf{y}}_k \quad \mathbf{b}_k \quad \dot{\mathbf{b}}_k]$ where,

- x_k represents the x-axis position at time epoch k
- $\dot{x}_k$ represents the velocity for the x-axis position at time epoch k
- y_k represents the y-axis position at time epoch k
- $\dot{y}_k$ represents the velocity for the y-axis position at time epoch k
- b_k represents the clock drift at time instant k
- $\dot{b}_k$ represents the clock drift velocity at time instant k

We know that the Kalman Filter is a recursive estimator using extensively the Gaussian assumption. Here, in comparison to LSE, we have a state transition model along with the measurement model. In the case of linear models, these are expressed as: [9]

(a) State Transition Model:

$$\mathbf{x}_{k+1} = \mathbf{F}_k \mathbf{x}_k + \mathbf{G}_k \mathbf{u}_k + \mathbf{v}_k \quad (3.9)$$

where,

- $\mathbf{F}_k$ is the assumed state transition matrix, providing the vehicle dynamics equations
- $\mathbf{u}_k$ is the command vector (known) that signifies the pilot/driver's actions
- $\mathbf{G}_k$ is the command matrix (known) providing a link between the pilot/driver's actions and the state
- $\mathbf{v}_k$ is the state noise vector or process noise, also called the modelization error

(b) Measurement Model:

$$\mathbf{y}_k = \mathbf{H}_k \mathbf{x}_k + \mathbf{n}_k \quad (3.10)$$

where,

- $\mathbf{x}_k$ is the state vector defined previously
- $\mathbf{y}_k$ is the measurement vector gathering all the pseudoranges at time k
- $\mathbf{n}_k$ is the measurement noise vector at time instant k that affects the measurement

Regarding the state model, in our study, we will use a first-order random walk model, defined by:

$$\mathbf{F}_k = \begin{bmatrix} 1 & \Delta t & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & \Delta t & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & \Delta t & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix} \quad (3.11)$$

Assumptions: The KF assumes that in equations 3.9 and 3.10, in addition to the linear aspect of these equations, the corresponding noises $\mathbf{v}_k$ and $\mathbf{n}_k$ are Gaussian. Therefore, both $\mathbf{v}_k$ and $\mathbf{w}_k$ are independent and Gaussian with distributions $\mathcal{N}(0, \mathbf{Q}_k)$ and $\mathcal{N}(0, \mathbf{R}_k)$ respectively. As Kalman Filter is a recursive process, an initial state must be provided that is Gaussian, namely $\mathbf{x}_0 \sim \mathcal{N}(\hat{\mathbf{x}}_0, \mathbf{\Sigma}_0)$. It can be shown that, due to the recursive equations, the system state vector $\mathbf{x}_k$ is also random and Gaussian, and is therefore described by two parameters, i.e., its mean ($\hat{\mathbf{x}}_k = \mathbf{E}[\mathbf{x}_k]$) and its covariance matrix ($\mathbf{\Sigma}_k = \text{Cov}[\mathbf{x}_k]$).

Principle: The principle of the Kalman Filter can be explained in the following two steps:

■ **State Prediction:** The first step in KF is the prediction phase where the internal state vector and the state covariance matrix are estimated based on only the state transition model. In other words, a prediction of the state distribution is obtained using the state transition model, and without the use of a new observation. As this distribution is Gaussian, it is fully determined by its mean and covariance matrix, expressed as

(a) **State prediction:**

$$\hat{\mathbf{x}}_{k|k-1} = \mathbf{F}_{k-1}\hat{\mathbf{x}}_{k-1|k-1} + \mathbf{G}_{k-1}\mathbf{u}_{k-1} \quad (3.12)$$

(b) **State Covariance prediction:**

$$\Sigma_{k|k-1} = \mathbf{F}_{k-1}\hat{\Sigma}_{k-1|k-1}\mathbf{F}_{k-1}^T + \mathbf{Q}_{k-1} \quad (3.13)$$

After the prediction step, a priori state is obtained given by: $\hat{\mathbf{x}}_{k|k-1} = \mathbf{E}[\mathbf{x}_k|\mathbf{y}_0, \dots, \mathbf{y}_{k-1}] = \mathbf{E}[\mathbf{x}_k|\mathbf{y}_0^{k-1}]$ with covariance matrix, $\Sigma_{k|k-1} = \mathbf{Cov}[\mathbf{x}_k|\mathbf{y}_0^{k-1}]$.

■ **Measurement Update:** The second step is the update step of the apriori estimated state vector and the covariance matrix at the current epoch k . Now, the predicted state distribution is used as the prior distribution with the measurement in order to obtain the Bayesian estimator of the state vector. In the linear Gaussian model, it can be shown that this distribution is also Gaussian, whose mean is denoted as $\hat{\mathbf{x}}_{k|k}$ and whose covariance matrix as $\Sigma_{k|k}$, which are expressed as

(a) **State Update:**

$$\hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \mathbf{K}_k(\mathbf{y}_k - \mathbf{H}_k\hat{\mathbf{x}}_{k|k-1}) \quad (3.14)$$

$\mathbf{K}_k$ denotes the Kalman gain at epoch k , as seen in equation 3.14, can be interpreted as the ratio between the estimated apriori covariance matrix $\Sigma_{k|k-1}$ and the sum of the estimated apriori covariance matrix and the observation covariance matrix $\mathbf{R}_k$ and is expressed as

$$\mathbf{K}_k = \Sigma_{k|k-1}\mathbf{H}_k^T(\mathbf{H}_k\Sigma_{k|k-1}\mathbf{H}_k^T + \mathbf{R}_k)^{-1} \quad (3.15)$$

(b) **State Covariance Update:**

$$\Sigma_{k|k} = (\mathbf{I} - \mathbf{K}_k\mathbf{H}_k)\hat{\Sigma}_{k|k-1} \quad (3.16)$$

After taking into account the k^{th} measurement $\mathbf{y}_k$, we obtain the posterior state given by: $\hat{\mathbf{x}}_{k|k} = \mathbf{E}[\mathbf{x}_k|\mathbf{y}_0, \dots, \mathbf{y}_k] = \mathbf{E}[\mathbf{x}_k|\mathbf{y}_0^k]$ with covariance matrix, $\Sigma_{k|k} = \mathbf{Cov}[\mathbf{x}_k|\mathbf{y}_0^k]$.

Kalman filter works with linear and Gaussian systems but our model is non-linear. We, therefore, resort to Extended Kalman filters that can deal with non-linear and Gaussian systems.

3.1.3 Extended Kalman Filter (EKF)

For nonlinear models, the Extended Kalman Filter (EKF) can be used. Kalman filters also assume Gaussian noise. The model equations are linearized around the previous estimate to be able to apply the Kalman Filter algorithm. This implies that now we have: $\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k) + \mathbf{G}_k\mathbf{u}_k + \mathbf{v}_k$ and $\mathbf{y}_k = \mathbf{h}(\mathbf{x}_k) + \mathbf{n}_k$, where $\mathbf{f}(\mathbf{x})$ and $\mathbf{h}(\mathbf{x})$ are non-linear functions. Therefore, now, we need to linearize the model. For the state transition model, the linearization point is the last estimated state, $\hat{\mathbf{x}}_{k|k}$, given by:

$$\mathbf{f}(\mathbf{x}_k) \approx \mathbf{f}(\hat{\mathbf{x}}_{k|k}) + \frac{\partial \mathbf{f}}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k})(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}). \quad (3.17)$$

Similarly, for the measurement model, the linearization point is the last predicted state, $\hat{\mathbf{x}}_{k|k-1}$, given by

$$\mathbf{h}(\mathbf{x}_k) \approx \mathbf{h}(\hat{\mathbf{x}}_{k|k-1}) + \frac{\partial \mathbf{h}}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1})(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1}). \quad (3.18)$$

Linearization: For linearizing the equations, first order Taylor's series expansion has been used and the system of equations becomes:

$$\mathbf{x}_{k+1} = \mathbf{f}(\mathbf{x}_k) + \mathbf{G}_k \mathbf{u}_k + \mathbf{v}_k = \mathbf{f}(\hat{\mathbf{x}}_{k|k}) + \mathbf{F}_k(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k}) + \mathbf{G}_k \mathbf{u}_k + \mathbf{v}_k \quad (3.19)$$

$$\mathbf{y}_k = \mathbf{h}(\mathbf{x}_k) + \mathbf{n}_k = \mathbf{h}(\hat{\mathbf{x}}_{k|k-1}) + \mathbf{H}_k(\mathbf{x}_k - \hat{\mathbf{x}}_{k|k-1}) + \mathbf{n}_k \quad (3.20)$$

where $\frac{\partial \mathbf{f}}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k}) = \mathbf{F}_k$ and $\frac{\partial \mathbf{h}}{\partial \mathbf{x}_k}(\hat{\mathbf{x}}_{k|k-1}) = \mathbf{H}_k$. Denoting the known quantities (the models and the Jacobians that are known analytically) as:

$$\tilde{\mathbf{u}}_k = \mathbf{G}_k \mathbf{u}_k + \mathbf{f}(\hat{\mathbf{x}}_{k|k}) - \mathbf{F}_k(\hat{\mathbf{x}}_{k|k}) \quad (3.21)$$

$$\tilde{\mathbf{y}}_k = \mathbf{h}(\hat{\mathbf{x}}_{k|k-1}) - \mathbf{H}_k \hat{\mathbf{x}}_{k|k-1} \quad (3.22)$$

Thus, we have:

$$\mathbf{x}_{k+1} = \mathbf{F}_k \mathbf{x}_k + \tilde{\mathbf{u}}_k + \mathbf{v}_k \quad (3.23)$$

$$\mathbf{y}_k - \tilde{\mathbf{y}}_k = \mathbf{H}_k \mathbf{x}_k + \mathbf{n}_k \quad (3.24)$$

Thus, this new system is now linear and the conventional KF algorithm can be used. Hence, the EKF algorithm can be written as:

■ **State Prediction:** $\hat{\mathbf{x}}_{k|k-1} = \mathbf{f}(\hat{\mathbf{x}}_{k-1|k-1}) + \tilde{\mathbf{u}}_{k-1}$

■ **State Covariance Prediction:** $\Sigma_{k|k-1} = \mathbf{F}_{k-1} \Sigma_{k-1|k-1} \mathbf{F}_{k-1}^T + \mathbf{Q}_{k-1}$

■ **Kalman Gain Computation:** $\mathbf{K}_k = \Sigma_{k|k-1} \mathbf{H}_k^T (\mathbf{H}_k \Sigma_{k|k-1} \mathbf{H}_k^T + \mathbf{R}_k)^{-1}$

■ **State Update:** $\hat{\mathbf{x}}_{k|k} = \hat{\mathbf{x}}_{k|k-1} + \mathbf{K}_k (\mathbf{y}_k - \mathbf{h}(\hat{\mathbf{x}}_{k|k-1}))$

■ **State Covariance Update:** $\Sigma_k = (\mathbf{I} - \mathbf{K}_k \mathbf{H}_k) \Sigma_{k|k-1}$

3.1.4 Particle Filter (PF)

Particle filters [6] are Bayesian filters suited for nonlinear and non-Gaussian models. They involve the implementation of Sequential Importance Sampling (SIS) which is a Monte-Carlo (MC) method. SIS is a technique to implement a recursive Bayesian filter by MC simulations. Thus, the idea behind the particle filter is to discretize the initial distribution with particles with associated weights, and to propagate them through the prior and the likelihood, resulting in a discrete version of the a posteriori distribution. Then the maximum a posteriori (MAP) estimator can be used to select the best path that fits the observations. In other words, the required posterior density function is represented by a set of random samples with associated weights and computes estimates based on these weights. As the number of samples becomes very large, the Monte-Carlo characterization becomes equivalent to the usual functional description of the posterior PDF, and the SIS filter approaches the optimal Bayesian estimate.

Algorithm: The estimation of the posterior density of the state variable is the main aim of the PF algorithm. In other words, the idea behind PF is to find out the sequential estimation of the values of the hidden states, $\mathbf{X}_k$, when the values of the observation process at any time k are given, i.e., $\mathbf{y}_0, \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_k$. The same has been expressed in figure 3.1.

As we know that a posterior density, $p(\mathbf{x}_k | \mathbf{y}_0, \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_k)$ is followed by all the Bayesian estimates of $\mathbf{X}_k$. An approximation of these conditional probabilities is provided by the PF technique. In contrast to it, the full posterior, $p(\mathbf{x}_0, \mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_k | \mathbf{y}_0, \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_k)$ is processed by the Markov Chain Monte-Carlo (MCMC) method. Thus, the model used for the particle filter is given as:

$$\mathbf{X}_k | \mathbf{X}_{k-1} = \mathbf{x}_k \sim p(\mathbf{x}_k | \mathbf{x}_{k-1}) \quad (3.25)$$

$$\mathbf{Y}_k | \mathbf{X}_k = \mathbf{y}_k \sim p(\mathbf{y}_k | \mathbf{x}_k) \quad (3.26)$$

Now comes the step for the approximation of the probability density of the PF, $p(\mathbf{x}_k | \mathbf{y}_0, \mathbf{y}_1, \mathbf{y}_2, \dots, \mathbf{y}_k)$ by a weighted set of the samples of the total number of particles, referred to as N_p . Thus, a particle is represented as a possible value for the state vector along with the corresponding probability, which is the value of the posterior PDF at this point, i.e., a particle is represented as:

$$(\mathbf{x}_k^{(i)}, \mathbf{w}_k^{(i)}), \quad \forall i = 1, 2, \dots, N_p \quad (3.27)$$

The principle of importance sampling (IS) relies on the Monte Carlo integration. Assume that we have N independent samples $\mathbf{x}_1, \dots, \mathbf{x}_N$ drawn from a distribution with PDF $p(\mathbf{x})$. Then, the Monte-Carlo approximations states that for any function f , we have:

$$\int f(\mathbf{x})p(\mathbf{x})d\mathbf{x} = \mathbf{E}_{\mathbf{x} \sim p(\mathbf{x})}[f(\mathbf{x})] \approx \frac{1}{N} \sum_{i=1}^N f(\mathbf{x}_i). \quad (3.28)$$

This last expression is equivalent to considering a discretization of the pdf $p(\mathbf{x})$ such as:

$$p(\mathbf{x}) \approx \frac{1}{N} \sum_{i=1}^N \delta(\mathbf{x} - \mathbf{x}_i) \quad (3.29)$$

where $\delta(\mathbf{x})$ denotes the Dirac distribution. However in some situations, drawing samples from $p(\mathbf{x})$ might be impossible for practical reasons. The idea of IS is to draw samples $\mathbf{x}_1, \dots, \mathbf{x}_N$ from a distribution $q(\mathbf{x})$ which contains the support of the target distribution $p(\mathbf{x})$, referred to as the proposal distribution, from which we are able to draw samples. Noting that

$$\int f(\mathbf{x})p(\mathbf{x})d\mathbf{x} = \int f(\mathbf{x}) \frac{p(\mathbf{x})}{q(\mathbf{x})} q(\mathbf{x})d\mathbf{x} = \mathbf{E}_{\mathbf{x} \sim q(\mathbf{x})} \left[\frac{p(\mathbf{x})}{q(\mathbf{x})} f(\mathbf{x}) \right] \quad (3.30)$$

and denoting

$$\mathbf{w}_i = \frac{p(\mathbf{x}_i)}{q(\mathbf{x}_i)} \quad (3.31)$$

the so-called importance weights, we have:

$$\int f(\mathbf{x})p(\mathbf{x})d\mathbf{x} = \mathbf{E}_{\mathbf{x} \sim p(\mathbf{x})}[f(\mathbf{x})] \approx \frac{1}{N} \sum_{i=1}^N \mathbf{w}_i f(\mathbf{x}_i). \quad (3.32)$$

This last expression is equivalent to considering a discretization of the PDF $p(\mathbf{x})$ such as:

$$p(\mathbf{x}) \approx \frac{1}{N} \sum_{i=1}^N \mathbf{w}_i \delta(\mathbf{x} - \mathbf{x}_i). \quad (3.33)$$

To improve this approximation, we rather use:

$$p(\mathbf{x}) \approx \frac{1}{N} \sum_{i=1}^N \bar{\mathbf{w}}_i \delta(\mathbf{x} - \mathbf{x}_i) \quad (3.34)$$

where

$$\bar{\mathbf{w}}_i = \frac{\mathbf{w}_i}{\sum_{i=1}^N \mathbf{w}_i} \quad (3.35)$$

to ensure that the approximate pdf has a unit integral. In the case of PF, the posterior distribution is therefore approximated as ...

One key element of PF is the proposal distribution: $\Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k})$. The target distribution is the optimal proposal distribution which would be the optimal choice, where the target distribution is given as :

$$\Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k}) = p(\mathbf{x}_k|\mathbf{x}_{k-1}, \mathbf{y}_k) = \frac{p(\mathbf{y}_k|\mathbf{x}_k)}{\int p(\mathbf{y}_k|\mathbf{x}_k)p(\mathbf{y}_k|\mathbf{x}_{k-1}d\mathbf{x}_k)} \cdot p(\mathbf{y}_k|\mathbf{x}_{k-1}d\mathbf{x}_k) \quad (3.36)$$

When $N_p \rightarrow \infty$, the proposal becomes: [5]

$$\Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k}) \approx \frac{p(\mathbf{y}_k|\mathbf{x}_k)}{\int p(\mathbf{y}_k|\mathbf{x}_k)\hat{p}(d\mathbf{x}_k|\mathbf{x}_{k-1})} \cdot \hat{p}(d\mathbf{x}_k) = \sum_{i=1}^{N_p} \frac{p(\mathbf{y}_k|\mathbf{x}_k^i(\mathbf{x}_{k-1}))}{\sum_{j=1}^{N_p} p(\mathbf{y}_k|\mathbf{x}_k^j(\mathbf{x}_{k-1}))} \cdot \delta_{\mathbf{x}_k^i(\mathbf{x}_{k-1})} \cdot d\mathbf{x}_k \quad (3.37)$$

Therefore, when we approximate the result, we have:

$$\hat{p}(d\mathbf{x}_k|\mathbf{x}_{k-1}) = \frac{1}{N_p} \sum_{i=1}^{N_p} \delta_{\mathbf{x}_k^i}(\mathbf{x}_{k-1}) \cdot d\mathbf{x}_k \approx_{N_p \uparrow \infty} p(\mathbf{x}_k|\mathbf{x}_{k-1})d\mathbf{x}_k \quad (3.38)$$

where we use the notation $\delta_{\mathbf{x}_k^i}(\mathbf{x}_{k-1}) = \delta(\mathbf{x}_{k-1} - \mathbf{x}_k^i)$. These are associated with N_p random samples, $\mathbf{x}_k^i(\mathbf{x}_{k-1})$, $i = 1, 2, 3, \dots, N_p$ with a conditional distribution of the random state, $\mathbf{x}_k$, given by $\mathbf{x}_k = \mathbf{x}_{k-1}$.

When the **importance density** is chosen to be the **transition prior probability**, it becomes easier to draw samples or particles and compute the corresponding importance weight.

$$\Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k}) = p(\mathbf{x}_k|\mathbf{x}_{k-1}) \quad (3.39)$$

where $\Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k})$ is the proposal density and $p(\mathbf{x}_k|\mathbf{x}_{k-1})$ is the prior density. With the SIS particle filters, there exists a problem of degeneracy. It is seen that after a few iterations, the weight will concentrate on a few particles only and most particles will have negligible weight. Degeneracy occurs when the weight of one particle is close to one and the weights of all other particles are close to or nearly zero. The filtering distribution in this case is represented by a single particle, which results in a very poor approximation. To solve this problem, the idea is to perform resampling because we want to eliminate particles that have small weights and concentrate on the particles with large weights, and hence, what we implement is the SIR (Sequential Importance Resampling) algorithm particle filter. The algorithm for a SIR PF is given by: [5], [6]

■ **Initialization:** The particles are randomly generated from the initial prior Gaussian distribution, $p(\mathbf{x}_k|\mathbf{x}_{k-1}, \mathbf{y}_k)$

■ For $i = 1, 2, \dots, N_p$, first the samples are drawn from the proposal: $\mathbf{x}_k^i \sim \Pi(\mathbf{x}_k|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k})$

■ For $i = 1, 2, \dots, N_p$, the importance weights are updated as the particle filters are also an approximation, therefore the non-linear filter equation requires recursive implementation, which is the update step: $\hat{w}_k^i = \hat{w}_{k-1}^i \cdot \frac{p(\mathbf{y}_k|\mathbf{x}_k^i)p(\mathbf{x}_k^i|\mathbf{x}_{k-1}^i)}{\Pi(\mathbf{x}_k^i|\mathbf{x}_{0:k-1}, \mathbf{y}_{0:k})}$.

Since the proposal distribution is chosen as the prior distribution, the above equation is reduced to:

$$w_k^i = w_{k-1}^i p(\mathbf{y}_k|\mathbf{x}_k^i) \quad (3.40)$$

■ The normalized importance weights are expressed as: $w_k^i = \frac{\hat{w}_k^i}{\sum_{j=1}^{N_p} \hat{w}_k^j}$

■ Now the computation of the estimate of the effective number of particles is performed as:

$$\hat{N}_{eff} = \frac{1}{\sum_{i=1}^{N_p} (w_k^i)^2}$$

■ If $\hat{N}_{eff} < N_{threshold}$, then resampling is performed, given by:

- (a) Draw N_p particles from the current particle set with probabilities proportional to their weights. Then replace the current particle set with this one.
- (b) For $i = 1, 2, \dots, N_p$; set $w_k^i = \frac{1}{N_p}$

3.2 Map matching

This section introduces the concept of map-matching and the methods used for its implementation with particle filters for GNSS positioning.

Map matching determines a vehicle's position on a digital road database. In this project, the map-matching technique is used along with GNSS observations in order to have a more accurate estimation of the vehicle's position. The map database provides a priori information for the system and acts like 'assistance' data. The advantages of implementing map matching involve cost-effective and reliable solution that provides confidence, and more accuracy solves data efficiency problems, and can overcome bad GNSS observations due to harsh environments.

The maps are defined by the line segments with nodes at the crossing. This technique uses non-standard priors. For example, at a crossing, the prior distribution is a mix of distributions containing all the possible directions that the vehicle can take. Particle filters can only be used in this case because classical filters cannot handle such non-standard priors and likelihood. In this project, two map-matching techniques were implemented with particle filters and a T-shaped map has been defined with a particular length.

3.2.1 Method 1: Map information is used in the observations

The first method is a simple process where we find the distance between the map and the generated particles at that epoch. These distances are combined inside the observation equation along with PVT equations. This technique is used where the observation equations are updated (or in other words, the likelihood PDF is changed) with no change in the proposal distribution. The following figure 3.2 shows the representation of the method used and the algorithm to compute the distance:

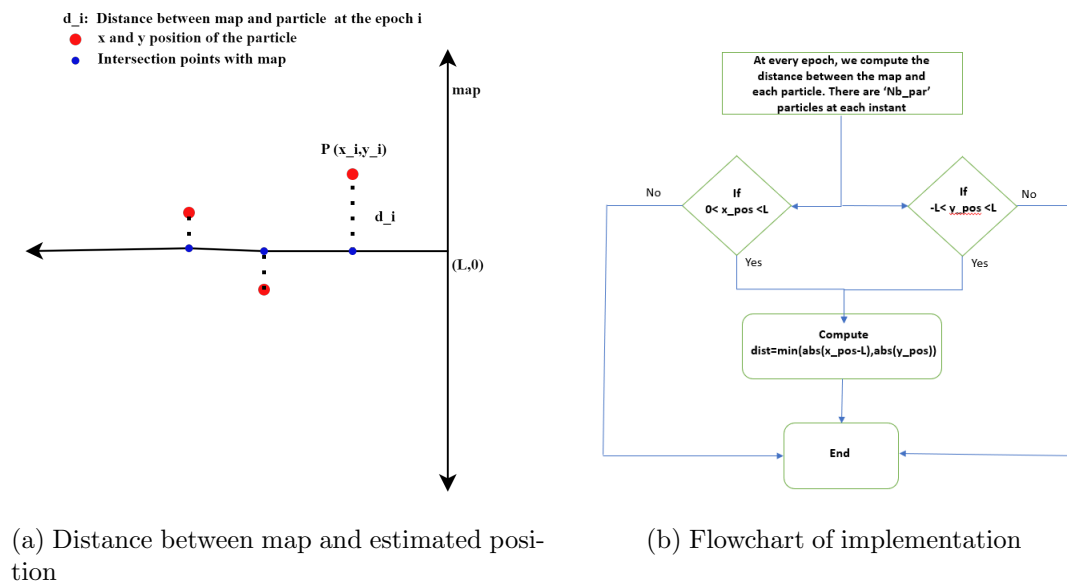


Figure 3.2: Method 1: Observations with the distance between map and particles

To account for the map, we are able at any time to measure the distance from a state vector $\mathbf{x}$ to the map, that we denote $d(\mathbf{x})$. If we assume that the relation between a position (x, y) and the closest point to the map (x_m, y_m) is as

$$\begin{bmatrix} x \\ y \end{bmatrix} = \begin{bmatrix} x_m \\ y_m \end{bmatrix} + \mathbf{e}_m \quad (3.41)$$

where $\mathbf{e}_m$ accounts for the uncertainty of the map and is as $\mathbf{e}_m \sim \mathcal{N}(0, \sigma_m^2 I_2)$, then

$$d(\mathbf{x}) \sim \mathcal{R}(\sigma_m^2) \quad (3.42)$$

where $\mathcal{R}(\sigma_m^2)$ denotes the Rayleigh distribution with parameter σ_m^2 whose support is $\mathbb{R}^+$ and whose PDF is

$$p(d) = \frac{d}{\sigma_m^2} \exp\left(-\frac{1}{2\sigma_m^2} d^2\right) \quad (3.43)$$

. We therefore augment the observations $\mathbf{y}_k$ as $\mathbf{z}_k = [\mathbf{y}_k^T \ d(\mathbf{x}_k)]^T$. Assuming the observations $\mathbf{y}_k$ and $d(\mathbf{x}_k)$ are independent, we have

$$p(\mathbf{z}_k | \mathbf{x}_k) = p(\mathbf{y}_k | \mathbf{x}_k) p(d(\mathbf{x}_k)). \quad (3.44)$$

Hence, the update (3.40) is as

$$\begin{aligned} w_k^i &= w_{k-1}^i p(\mathbf{y}_k | \mathbf{x}_k^i) p(d(\mathbf{x}_k^i)) \\ &\propto w_{k-1}^i \exp\left(-\frac{1}{2} [\mathbf{y}_k - h(\mathbf{x}_k^i)]^T \mathbf{R}_k^{-1} [\mathbf{y}_k - h(\mathbf{x}_k^i)]\right) d(\mathbf{x}_k^i) \exp\left(-\frac{d(\mathbf{x}_k^i)^2}{2\sigma_m^2}\right). \end{aligned} \quad (3.45)$$

3.2.2 Method 2: Proposal drawn from possible position of the vehicle on the map

The second method involves defining the proposal distribution of the particle filter with the possible directions or positions that a vehicle can go from an estimated position, at an epoch. To implement this, we used the method of the intersection of circles. At each epoch, a circle is drawn with the center as the position estimated by PF at the previous epoch. The radius of the circle (r_i) is defined as $r_i = \sqrt{v_x^2 + v_y^2} \delta t$; where v_x and v_y are the estimated x and y velocity at the previous epoch. The intersection points of the circle with the map are computed. We also compute the number of intersection points at every epoch. A uniform proposal distribution is drawn from the intersection points. The points are assumed to have equal probabilities. For the computation of the intersection points, we consider a T-shaped map. The map can be seen to be defined by two straight-line segments. The horizontal line segment lies on the x-axis with $x \in [0, L]$ and $y=0$. The vertical line segment is parallel to the y-axis with $y \in [-L, L]$ and $x=L$, where L is the length of the map. The equations for the computation of points of intersection (x_{int}, y_{int}) are:

Horizontal line-segment

$y=0$ and $0 \leq x \leq L$

Let the position estimate at the previous epoch be x_i, y_i . We first check if the condition: $r_i^2 > y_i^2$ is true. Then the points of intersection are

$$x_{int} = x_i \pm \sqrt{r_i^2 - y_i^2} \quad (3.46)$$

and $y_{int} = 0$.

Vertical line-segment

$x=L$ and $-L \leq y \leq L$

Let the position estimate at the previous epoch be x_i, y_i . We first check if the condition: $r_i^2 > (L - x_i)^2$ is true. Then the points of intersection are

$$y_{int} = x_i \pm \sqrt{r_i^2 - (L - x_i)^2} \quad (3.47)$$

and $x_{int} = L$.

Once we have this information, certain equations for the particle filter are changed to include map matching like the proposal and weight update equations, which are implemented in our algorithm.

The following figure 3.3 shows the circles whose intersections are to be found at each epoch and the algorithm to find them:

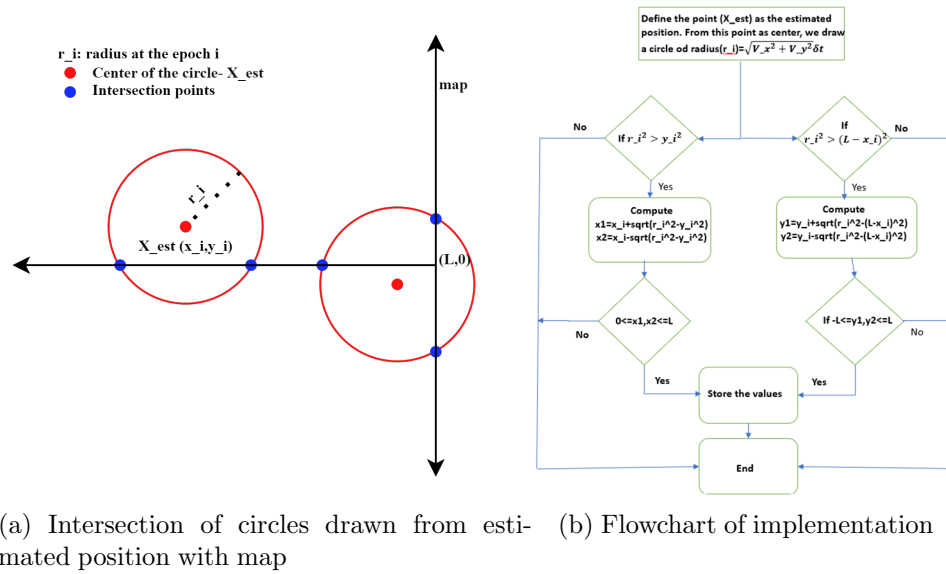


Figure 3.3: Method 2: New proposal density

The equations that need to be updated are:

■ The proposal density:

$$\Pi(x_k | x_{0:k-1}, y_{0:k}) = \frac{1}{N_{int}} \sum_{j=1}^{N_{int}} \delta(x_k^i - X_{int}^i) \quad (3.48)$$

where, N_{int} is the number of intersections, x_k is the estimate from PF and X_{int} is the points of intersection of the map with the circles.

■ For $i=1, 2, \dots, N_p$; samples are drawn from this proposal:

$$x_k^i \sim \Pi(x_k | x_{0:k-1}, y_{0:k}) = \frac{1}{N_{int}} \sum_{j=1}^{N_{int}} \delta(x_k^i - X_{int}^i) \quad (3.49)$$

■ For $i=1, 2, \dots, N_p$; the importance weights are updated:

$$\hat{w}_k^i = \hat{w}_{k-1}^i \frac{p(y_k | x_k^i) p(x_k^i | x_{k-1}^i)}{\Pi(x_k | x_{0:k-1}, y_{0:k})} = \hat{w}_{k-1}^i \frac{p(y_k | x_k^i) p(x_k^i | x_{k-1}^i)}{\frac{1}{N_{int}} \sum_{j=1}^{N_{int}} \delta(x_k^i - X_{int}^i)} \quad (3.50)$$

Since for each intersecting point, $\sum_{j=1}^{N_{int}} \delta(x_k^i - X_{int}^i) = 1$; we have reduced equations for weight update:

$$\hat{w}_k^i = \hat{w}_{k-1}^i \frac{p(y_k|x_k^i)p(x_k^i|x_{k-1}^i)}{\frac{1}{N_{int}}} \quad (3.51)$$

The prior density, normalization of weights, and effective number of particle computations for resampling remain the same as in the SIR particle filter. The resampling of particles is done from the proposal density.