

Statistical Signal Processing
Prof. Prabin Kumar Bora
Department of Electronics and Electrical Engineering
Indian Institute of Technology – Guwahati

Lecture – 36
Vector Kalman Filter

(Refer Slide Time: 00:44)

Let us recall

- ❖ The scalar Kalman filter uses the first-order AR signal model
$$X(n) = aX(n-1) + W(n)$$

and the observation model

$$Y(n) = X(n) + V(n)$$

- ❖ Kalman filter recursively estimates the signal.
- ❖ We derived the filter equations through the innovation representation of $Y(n)$.

Hello students welcome to this lecture on vector Kalman filter. Let us recall the scalar Kalman filter uses the first order autoregressive signal model that is $X_n = a \text{ times } X_{n-1} + W_n$ where W_n is a white noise and is a constant to tackle non-stationarity. It can be function of n the observation model is given by $Y_n = X_n + V_n$. Again V_n is a white noise uncorrelated with X_n and uncorrelated with W_n .

Kalman filter recursively estimates the signal. So, that we know that Kalman filter recursively estimate the signal on the basis of the current and all the previous data. We derived the filter equations true innovation representation of Y_n for data Y_n we had a innovation representation and then we derived the scalar Kalman filter.

(Refer Slide Time: 01:48)

Let us recall....

The following are the scalar Kalman filter steps

-Calculate MSPE

$$P(n|n-1) = a^2 P(n-1|n-1) + \sigma_w^2$$

-Calculate Kalman gain

$$k_n^{(n)} = \frac{P(n|n-1)}{P(n|n-1) + \sigma_v^2}$$

-Estimate signal

$$\hat{X}(n|n) = a\hat{X}(n-1|n-1) + k_n^{(n)}(Y(n) - a\hat{X}(n-1|n-1))$$

-Update MS estimation error

$$P(n|n) = (1 - k_n^{(n)})P(n|n-1)$$

❖ We will extend the scalar Kalman filter to the vector case. The analysis is similar. We will state the salient results only.

The following are the scalar Kalman filtering steps. First calculate the mean square prediction error this formula P of n given $n - 1 = a^2$ times P of $n - 1$ given $n - 1 + \sigma_w^2$ so this is the mean square error estimated at time $n - 1$ then we calculate the Kalman gain K_n that equal to P of n given $n - 1$ that is the mean square prediction error divided by P of n given $n - 1$ means square prediction error + σ_v^2 .

Then we estimate the signal X of n given $n = a$ times X at $n - 1$ given $n - 1$ this is the predicted part and there is a correction part here into K_n that is Kalman gain multiplied by the innovation part $Y_n - a$ times X at $n - 1$ given $n - 1$ this is the estimator at instant $n - 1$. Then we update the MS estimation error P of n given $n = 1 - K$ and n into P of n given $n - 1$. This is the step for updating the mean square error now we will extend this scalar Kalman filter to the vector case.

The analysis is similar we will state the salient results only we will not go into detail analysis, but this derivation is also based on the innovation representation of the data. Now we will introduce vector Kalman filter.

(Refer Slide Time: 03:53)

Vector Kalman Filter

- As we have seen in the earlier, an AR(p) signal can be represented in a state-space model.
- A state-space model in general represents a linear time varying system and comprises two matrix equations: a state equation and an observation equation.
- The vector Kalman filter is a recursive estimator of the states of such models in which states are driven by noise and observations are made in presence of noise.
- There are number of ways to derive the Kalman filter. We will present the approach of LMMSE estimation and outline the derivation through the innovation representation of a WSS signal.

As we have seen in the earlier an ARP signal can be represented in a state space model. A state-space model in general represents a linear time varying system and comprises two matrix equations a state equation and an observation equation. This is the state space model used by the control engineers. The vector Kalman filter is a recursive estimator of the states of such model in which states are driven by noise we will solve the model is the state driven by noise and observations are made in the presence of noise.

So, that way we have noisy observations and we have to estimate the states. There are number of ways to derive the Kalman filter for example there is a Bayesian approach. We will present the applause of linear mean square error estimation and outline the derivation through the innovation representation of a WSS signal.

(Refer Slide Time: 05:04)

State-space Model

- ❖ This model arises from the p -th order difference equation representing signal or a linear system. *It comprises two matrix equations:*
- The state equation is
$$\mathbf{X}(n) = \mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n)$$
 first-order difference equation
 where $\mathbf{X}(n) = [X_1(n), X_2(n), \dots, X_p(n)]^T$ is the $p \times 1$ state vector,
 $\mathbf{A}(n)$ is a $p \times p$ system matrix and
 $\mathbf{W}(n)$ is the 0-mean noise vector with the $p \times p$ covariance matrix $\mathbf{Q}_p = E\mathbf{W}(n)\mathbf{W}^T(n)$.
- ❖ $\mathbf{W}(n)$ is assumed to be a white noise. Thus, $\mathbf{Q}_p$ is a diagonal matrix.
- ❖ In Bayesian derivation, $\mathbf{W}(n)$ is assumed to be jointly Gaussian.

Let us discuss this state space model this model arises from the p th order difference equation representing signal or a linear system. For example, in the case of ARP signal autoregressive signal of order P , we have a p th order difference equation. For example, in the case of ARP model we have a p th order difference equation. Similarly, a discrete-time linear system can be described by a p th order difference equation.

Now the state space model comprises of two situations it comprises two matrix equation actually matrix equations number one is state equation this state equation is given by this $\mathbf{X}_n = \mathbf{A}_n \mathbf{X}_{n-1} + \mathbf{W}_n$. So, here this is a first-order difference equation unlike in the case of general model where it was p th order differentiation this is first order difference equation. Now here $\mathbf{X}_n$ is the state vector comprises of P state variable x_1 and x_2 n up to X_p .

And so this is this state vector and $\mathbf{A}_n$ is a p/p because P dimensional state vector is there so $\mathbf{A}_n$ will be p/p system matrix or state matrix. $\mathbf{W}_n$ is a 0 mean noise vector with the p/p covariance matrix that is $\mathbf{Q}_w = E$ of $\mathbf{W}_n$ into $\mathbf{W}_n$ transpose. This $\mathbf{W}_n$ it is a noise vector which is known also known as the process noise and it is state component will have corresponding noise addition and this noise is a white noise.

And it is characterized by the covariance matrix E of $\mathbf{W}_n$ into $\mathbf{W}_n$ and transpose since $\mathbf{W}_n$ is assumed to be a white noise $\mathbf{Q}_W$ is a diagonal matrix. So this $\mathbf{Q}_W$ matrix is a diagonal matrix. In

Bayesian derivation W_n is assumed to be jointly Gaussian also in addition to white this is assumed to be Gaussian also.

(Refer Slide Time: 07:47)

State-space model...

- ❖ The observed data $\mathbf{Y}(n) = [Y_1(n) Y_2(n) \dots Y_q(n)]^T$ is related to the states by the observation equation $\mathbf{Y}(n) = \mathbf{c}(n)\mathbf{X}(n) + \mathbf{V}(n)$ where
- $\mathbf{c}(n)$ is a $q \times p$ output matrix and
- $\mathbf{V}(n)$ is the 0 mean white Gaussian noise vector with $q \times q$ covariance matrix

$$\mathbf{Q}_v = E\mathbf{V}(n)\mathbf{V}^T(n)$$

- ❖ $\mathbf{Q}_v$ is a diagonal matrix
- ❖ $\mathbf{V}_j(n)$ is uncorrelated with $\mathbf{X}(n)$ and $\mathbf{W}(n)$

Now we will see the observation model the observed data $\mathbf{Y}_n$ that is comprising of $Y_1 Y_2$ up to Y_q there are q observation at a instant n . So, this is represented in terms of vector and that is the $\mathbf{Y}_n$ vector it is related to the states by the observation equation that is $\mathbf{Y}_n = \mathbf{c}_n \mathbf{X}_n + \mathbf{V}_n$. So, this is the observation equation $\mathbf{Y}_n = \mathbf{c}_n \mathbf{X}_n + \mathbf{V}_n$. So this is the observation equation now here $\mathbf{c}_n$ is a q / p output matrix because this $\mathbf{Y}_n$ is a Q dimensional vector therefore this matrix is q/p matrix.

$\mathbf{V}_n$ is a 0 mean white Gaussian noise vector this is their measurement error or measurement noise which is also assumed to be white Gaussian is necessary only in the basic derivation and its covariance matrix is a q / q covariance matrix which is given by $\mathbf{Q}_v = E\mathbf{V}_n \mathbf{V}_n^T$ $\mathbf{Q}_v$ is a diagonal matrix because $\mathbf{V}_n$ is a white noise vector and further it is assumed that $\mathbf{V}_n$ is uncorrelated with both $\mathbf{X}_n$ and $\mathbf{W}_n$. So, this $\mathbf{V}_n$ is uncorrelated with $\mathbf{X}_n$ and also with the step noise that is $\mathbf{W}_n$. So, therefore $\mathbf{V}_n$ is uncorrelated with $\mathbf{X}_n$ that is the state and the $\mathbf{W}_n$ that is the state noise or process noise.

(Refer Slide Time: 09:47)

LMMSE estimation and LMMSE prediction

❖ Denote the LMMSE estimator

$$\hat{\mathbf{X}}(n|n) = \hat{E}(\mathbf{X}(n) | \mathbf{Y}(0), \mathbf{Y}(1), \dots, \mathbf{Y}(n)) \text{ and}$$

the LMMSE predictor

$$\hat{\mathbf{X}}(n|n-1) = \hat{E}(\mathbf{X}(n) | \mathbf{Y}(0), \mathbf{Y}(1), \dots, \mathbf{Y}(n-1)).$$

❖ $\hat{\mathbf{X}}(n|n-1)$ is a one step predictor based on all past data

$$\begin{aligned} \hat{\mathbf{X}}(n|n-1) &= \hat{E}(\mathbf{X}(n) | \mathbf{Y}(0), \mathbf{Y}(1), \dots, \mathbf{Y}(n-1)) \\ &= \hat{E}(\mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n) | \mathbf{Y}(0), \mathbf{Y}(1), \dots, \mathbf{Y}(n-1)) \end{aligned}$$

(Using the state equation)

$$\therefore \hat{\mathbf{X}}(n|n-1) = \mathbf{A}(n)(\hat{\mathbf{X}}(n-1|n-1))$$

Now let us see linear minimum mean square error estimation and linear minimum mean square error prediction associated with Kalman filter. Denote the LMMSE estimator by this $\hat{\mathbf{X}}(n|n)$ given n . $\hat{E}$ is the symbol for linear minimum mean square error estimator. Therefore, $\hat{\mathbf{X}}(n|n)$ given n is $\hat{E}$ of $\mathbf{X}(n)$ given all the data up to present $\mathbf{Y}(0), \mathbf{Y}(1)$ up to $\mathbf{Y}(n)$ and the linear minimum mean square error predictor is given by this $\hat{\mathbf{X}}(n|n-1)$ data up to $n-1$.

So, that way it is $\hat{E}$ of $\mathbf{X}(n)$ given $\mathbf{Y}(0), \mathbf{Y}(1)$, up to $\mathbf{Y}(n-1)$. So, therefore this is a one-step predictor based on all the past data. Now this quantity $\hat{\mathbf{X}}(n|n-1)$ that we can write as $\hat{E}$ of $\mathbf{X}(n)$ given $\mathbf{Y}(0), \mathbf{Y}(1)$ up to $\mathbf{Y}(n-1)$. Now $\mathbf{X}(n)$, we can write like $\mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n)$. So, this is the state model we are using so using this state equation we will get the above expression = $\hat{E}$ of $\mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n)$ given $\mathbf{Y}(0), \mathbf{Y}(1)$ up to $\mathbf{Y}(n-1)$ and now this is $\hat{E}$ of $\mathbf{X}(n-1)$ given $\mathbf{Y}(0), \mathbf{Y}(1)$ up to $\mathbf{Y}(n-1)$.

Let LMMSE or linear minimum mean square error estimation of $\mathbf{X}(n-1)$ given data up to $n-1$ so that way it is $\hat{\mathbf{X}}(n-1|n-1)$ therefore $\hat{\mathbf{X}}(n|n-1)$ is $\mathbf{A}(n)$ times $\hat{\mathbf{X}}(n-1|n-1)$. So, this prediction is related to the previous estimation by this relationship $\mathbf{A}(n)$ times the previous estimation.

(Refer Slide Time: 12:09)

Error covariance matrices

- ❖ The corresponding state estimation errors are:

$$\mathbf{e}(n|n) = \mathbf{X}(n) - \hat{\mathbf{X}}(n|n) \quad (\text{estimation error})$$

and

$$\mathbf{e}(n|n-1) = \mathbf{X}(n) - \hat{\mathbf{X}}(n|n-1) \quad (\text{prediction error})$$

- ❖ The estimation error covariance matrix

$$\mathbf{P}(n|n) = E\mathbf{e}(n|n)\mathbf{e}'(n|n)$$

- ❖ The prediction error covariance

$$\mathbf{P}(n|n-1) = E\mathbf{e}(n|n-1)\mathbf{e}'(n|n-1)$$

6

We can find out the corresponding error covariance matrices also suppose estimation error $\mathbf{e}$ of n given $n = \mathbf{X}_n - \hat{\mathbf{X}}_n$ given n this is the estimator of $\mathbf{X}_n$ given data up to n . Similarly $\mathbf{e}$ of n given $n - 1$ that is $\mathbf{X}_n$ minus the prediction; that is the prediction of $\mathbf{X}_n$ given data up to $n - 1$. So, that way this is the prediction error this is the estimation error. Now we can define the error covariance matrix in scalar case we are interested only in mean square prediction error but here we will be considering the error covariance matrix.

So, that way $\mathbf{P}$ of n given n that is the notation for error covariance matrix. So, this is the estimation error covariance matrix and it is given by E of $\mathbf{e}_n$ given n that is the estimation error into $\mathbf{e}_n$ given n transpose. So, estimation error vector into estimation error vector transpose. So, this is the covariance matrix so estimation error covariance matrix. Similarly, prediction error covariance matrix we can find out that is denoted by $\mathbf{P}$ of n given $n - 1$ this is expected value of $\mathbf{e}_n$ given $n - 1$ into $\mathbf{e}_n$ given $n - 1$ transpose. This is the prediction error vector.

This is the transpose of the prediction error. So, that way we get the two covariance matrices that is $\mathbf{P}$ of n given n this is the estimation error covariance matrix and another covariance matrix is prediction error covariance matrix so that is given by this $\mathbf{p}$ of n given $n - 1 = E$ of $\mathbf{e}_n$ given $n - 1$ into $\mathbf{e}_n$ given $n - 1$ transpose prediction error vector into prediction error vector transpose.

(Refer Slide Time: 14:33)

Recursive estimation of states

- ❖ We have the innovation process

$$\tilde{Y}(n) = Y(n) - c(n)A(n)\hat{X}(n-1|n-1)$$

- ❖ Generalising the equation in the scalar model, we have

$$\begin{aligned}\hat{X}(n|n) &= \hat{X}(n|n-1) + k_s(Y(n) - \hat{Y}(n|n-1)) \\ &= A(n)\hat{X}(n-1|n-1) + k_s(Y(n) - c(n)\hat{X}(n-1|n-1)) \\ &= A(n)\hat{X}(n-1|n-1) + k_s(Y(n) - c(n)A(n)\hat{X}(n-1|n-1))\end{aligned}$$

where k_s is the $p \times q$ Kalman gain matrix.

Thus,

$$\hat{X}(n|n) = A(n)\hat{X}(n-1|n-1) + k_s(Y(n) - c(n)A(n)\hat{X}(n-1|n-1))$$

$\hat{X}(n|n) = a\hat{X}(n-1|n-1) + k_n \text{ innovation}$

Now let us see the innovation representation for a vector case we know the innovation representation for this scalar case. So, the derivation of the Kalman filter is based on the innovation of Y_n and which is given by this expression $\tilde{Y}$ and $E = Y_n$ - the linear estimator of Y_n given data up to $n-1$. So, $\hat{e}$ of Y_n given Y_0, Y_1 up to Y_{n-1} ; we have Y_n ; this is the Y_n vector = c_n matrix into state vector X_n vector + V_n , V_n is a white noise vector. So, this is the observation model that we are substituting here.

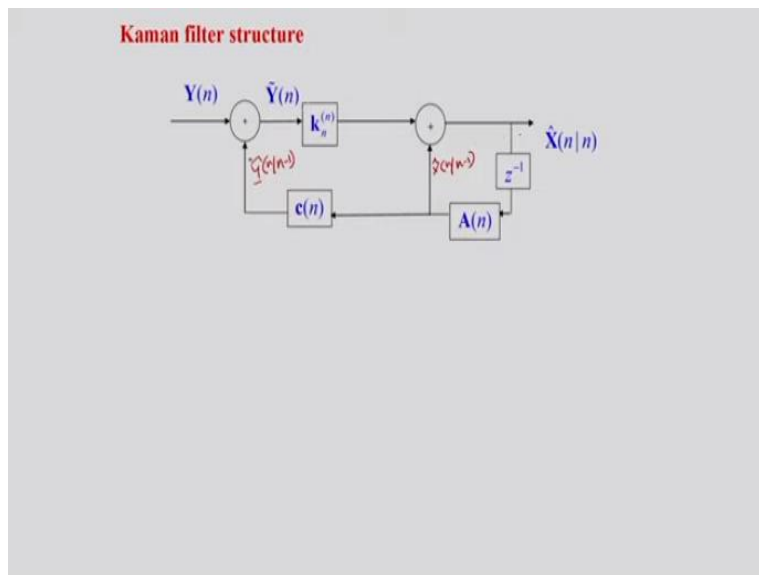
So, this will be E hat of C_n into $X_n + V_n$ given Y_0, Y_1 up to Y_{n-1} and this will be because now V_n this is the noise which is independent of all path data therefore this part will contribute 0 therefore we will have $Y_n - c_n$ into X hat n given $n-1$. Now this quantity also X hat n given $n-1 = A_n X$ hat $n-1$ given $n-1$. The prediction is obtained through the model that is A_n it into X hat $n-1$ given $n-1$ therefore $\tilde{Y}$ tilde n will be equal to $Y_n - c_n$, c_n matrix into A_n into X hat $n-1$ given $n-1$ this is the previous estimation.

Thus we have obtained the innovation of data that is given by $\tilde{Y}$ tilde $n = Y_n - c_n$ into A_n into X hat $n-1$ given $n-1$. So, this has come through the model we can use the innovation to get the recursive estimation of states we have $\tilde{y} = Y_n - c_n$ into A_n into X hat $n-1$ given $n-1$. Now generalizing this scalar model in scalar model we had X hat n given $n = A$ times X hat $n-1$ given $n-1 + K_n$ times K_n times the innovation and here we omit this super fixed n .

So, simply we write K_n therefore $\hat{X}_n$ in this case it will be a vector now $\hat{X}_n$ given $n =$ that predictor part that is $\hat{X}_n$ given $n - 1$ plus a correction based on the innovation that is K_n times $Y_n - \hat{Y}_n$ given $n - 1$ this is the prediction and this we can write as A_n into $\hat{X}_{n-1}$ given $n - 1$ that is the estimation part at the instant $n - 1$ into K_n times $Y_n - \hat{Y}_n$ and this we can write through the model c_n into $\hat{X}_n$ given $n - 1$.

Now we can write here also $\hat{X}_n$ given $n - 1 = A_n$ into $\hat{X}_{n-1}$ given $n - 1$. Therefore, our $\hat{X}_n$ given n will be n into $\hat{X}_{n-1}$ given $n - 1 + K_n$ times $Y_n - c_n$ into A_n into $\hat{X}_{n-1}$ given $n - 1$ where this K_n , K_n is the p/q Kalman gain matrix. This is the matrix ways the innovation part thus we can write $\hat{X}_n$ given $n = A_n$ into $\hat{X}_{n-1}$ given $n - 1 + K_n$ times $Y_n - c_n$ into A_n into $\hat{X}_{n-1}$ given $n - 1$. So this is the state estimator for the Kalman filter from this expression we can have the Kalman filter structure.

(Refer Slide Time: 19:56)



So, there is data here data is Y_n and the state estimator is $\hat{X}_n$ given n . Now this estimator is delayed by one unit and then multiplied by A_n to get the prediction $\hat{X}_n$ given $n - 1$. So, we will get here $\hat{X}_n$ given $n - 1$ given that up to $n - 1$ we will get the prediction below through the delayed version of this state estimator. So, when we will delay this, we will get $\hat{X}_{n-1}$ given $n - 1$ that will be multiplied by A_n and we will get this prediction.

And this result is multiplied by C_n matrix and then we will get the predicted $\hat{Y}_n$ that is $\hat{Y}$ at n given $n - 1$. So, this is the prediction of the data vector at instant n because of the data up to $n - 1$ so this is the prediction vector we can give a bar here to denote vector. So, that way this is the prediction vector this is the actual data and then difference is the innovation and that innovation is scaled by the Kalman gain and then edit with the predicted part to get the estimator. So, that way estimator has a prediction part here and then there is a correction part which is the innovation multiplied by the Kalman gain.

(Refer Slide Time: 21:43)

Kalman filter equations

- ❖ The estimation of the signal is based on the recursive estimation of the error covariance matrices and the Kalman gain
- ❖ Following the same procedure as the scalar Kalman filter, we can derive the following Kalman filter equations:
 - (1) Update of prediction error covariance matrix

$$P(n|n-1) = A(n)P(n-1|n-1)A'(n) + Q_w$$
 - (2) Estimating Kalman gain

$$K_n = P(n|n-1)C'(n)(C(n)P(n|n-1)C'(n) + Q_v)^{-1}$$
 - (3) Update of estimation error covariance matrix

$$P(n|n) = (I - K_n C(n))P(n|n-1)$$

Given the initial estimates $\hat{X}(-1|-1)$ and $P(-1|-1)$ the above equations can be used for recursive state estimation:

$$\hat{X}(n|n) = A(n)\hat{X}(n-1|n-1) + K_n(Y(n) - C(n)A(n)\hat{X}(n-1|n-1))$$

Now we will state the Kalman filter equation the estimation of the signal is based on the recursive estimation of the error covariance matrices and Kalman gain. So, we will see how they are estimated following this same procedure as the scalar Kalman filter we can derive the following Kalman filter equations. That is update of prediction error covariance matrix P of n given $n - 1 = n$ times P_{n-1} given $n - 1$ into A_n transpose. A_n is the time varying system matrix + Q_w , Q_w is the covariance matrix of the process noise.

So, therefore the prediction error covariance is given by this P of n given $n - 1 = A_n$ into P of $n - 1$ given $n - 1$ this is the estimation error covariance matrix multiplied by n transpose + Q_w . First one is the update of the prediction error covariance matrix and this is given by P of n given $n - 1 = A_n$ into P of $n - 1$ given $n - 1$ into A_n transpose + q_w where A_n is the system matrix or state matrix and QW is the covariance matrix of the noise vector.

So, that way we can find out the prediction error covariance in terms of the estimation error covariance of previous iteration. Then estimating the Kalman gain so for that we have $K_n = P_n C_n^T (C_n P_n C_n^T + Q_v)^{-1}$ of n given $n - 1$ we should obtain here into C_n transpose this is the output matrix transpose into C_n into P_n given $n - 1$ into C_n transpose + Q_v whole inverse. So, this part is a matrix and inverse of this is taken.

So, that way the Kalman gain is given by the error covariance matrix multiplied by this output matrix transpose multiplied by this inverse, inverse of C_n into P_n given $n - 1$ into C_n transpose + Q_v , Q_v is the covariance matrix of the other observation noise. So, once we have the Kalman gain we can update the estimation error covariance P of n given $n = I - K_n C_n$ this is the Kalman gain into C_n into P of n given $n - 1$.

So, given the prediction error covariance we can get the estimation error covariance. However along with this equation we must have the initial estimate that is $\hat{X}_{n-1}$ given $n - 1$ and the estimation for P of $n - 1$ given $n - 1$ that is the covariance matrix of the mean square error one estimate must be there. Now given these initial estimates $\hat{X}_{n-1}$ given $n - 1$ and P_{n-1} given $n - 1$ the above equations these 3 equations can be used to recursively estimate the state and the recursive state estimator is given by this $\hat{X}_n = A_n \hat{X}_{n-1} + K_n (Y_n - C_n A_n \hat{X}_{n-1})$. So, this is the innovation part, and this is scaled by K_n and added to the predicted part.

(Refer Slide Time: 26:14)

Vector Kalman filter algorithm

- ❖ State equation: $\mathbf{X}(n) = \mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n)$
- ❖ Observation equation: $\mathbf{Y}(n) = \mathbf{C}(n)\mathbf{X}(n) + \mathbf{V}(n)$
- ❖ Given (a) State matrix $\mathbf{A}(n), n = 0, 1, 2, \dots$ and the process noise covariance matrix $\mathbf{Q}_w$
(b) Observation parameter matrix $\mathbf{C}(n), n = 0, 1, 2, \dots$ and the observation noise covariance matrix $\mathbf{Q}_v$
(c) Observed data $\mathbf{Y}(n), n = 0, 1, 2, \dots$
- Initialization:
 $\hat{\mathbf{X}}(-1|-1) = 0$
 $\mathbf{P}(-1|-1) = E\mathbf{X}(-1)\mathbf{X}^T(-1) = \sigma^2 \mathbf{I}$

With these backgrounds we can now state the Kalman filter by algorithm. So, here we have to know this state equation $\mathbf{X}_n = \mathbf{A}_n \mathbf{X}_{n-1} + \mathbf{W}_n$ observation equation that is given by $\mathbf{Y}_n = \mathbf{C}_n \mathbf{X}_n + \mathbf{V}_n$ this is the observation equation. So, we must be given a state matrix $\mathbf{A}_n$ for $n = 0$ to P because it is p th order system or signal, and the process noise covariance matrix is given by $\mathbf{Q}_w$ which is a diagonal matrix.

The observation parameter matrix that is $\mathbf{C}_n$ for $n = 0, 1$ up to 2 if it is time varying otherwise it will be a simply $\mathbf{C}$ matrix and the observation noise covariance matrix is given by $\mathbf{Q}_v$ this is again a diagonal matrix and observed data $\mathbf{Y}_n$ that is from $n = 0, 1, 2$ etc. So, these are the things needed for the algorithm now it will be initialized as $\hat{\mathbf{X}}(-1|-1) = 0$ all states are initialized with 0 and this estimation error covariance matrix at instant -1 .

There is no estimation therefore this will be the covariance of the data itself people would know this statistical model we can use that and here we can write it as $\sigma^2 \mathbf{I}$. This is the estimator and if we take this type of initialization $\sigma^2 \mathbf{I}$ where $\mathbf{I}$ is an identity matrix then the Kalman equations converges. So, that way this σ^2 is a positive number which can be based on the estimation of this quantity or we can assume some value.

(Refer Slide Time: 28:25)

Vector Kalman filter algorithm ...

- Estimation:

For $n = 0, 1, 2, \dots$ do

Prediction

-Predict the state

$$\hat{X}(n|n-1) = A(n)\hat{X}(n-1|n-1)$$

-Estimate *a priori* error covariance matrix

$$P(n|n-1) = A(n)P(n-1|n-1)A'(n) + Q_w$$

Compute Kalman gain

$$k_n = P(n|n-1)c'(n)(c(n)P(n|n-1)c'(n) + Q_v)^{-1}$$

$$\hat{X}(n|n) = \hat{X}(n|n-1) + k_n(Y(n) - c(n)\hat{X}(n|n-1))$$

Update *a posteriori* error covariance matrix

$$P(n|n) = (I - k_n c(n))P(n|n-1)$$

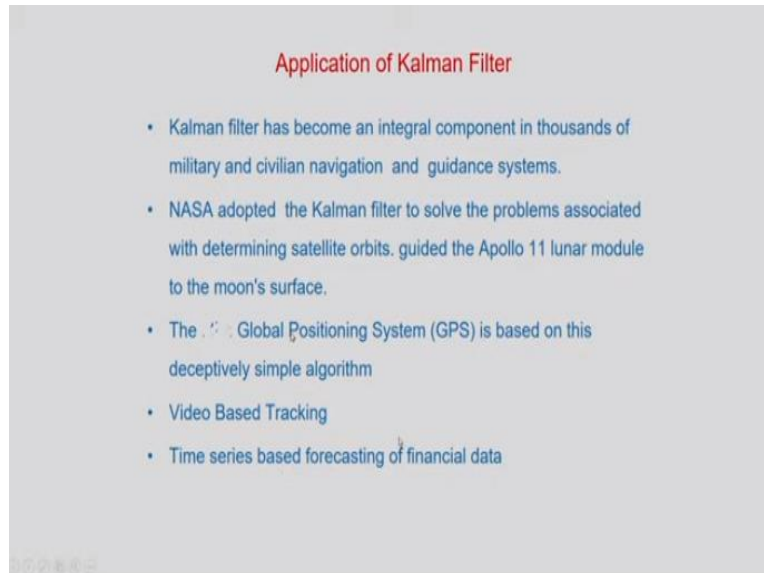
Kalman filter runs like this for $n = 0, 1, 2$ etc first we will predict so predict this state that is $\hat{X}(n|n-1)$ that is given by $A(n)$ into $\hat{X}(n-1|n-1)$. So this is the estimated value this predicted term will be used in the updating here then we will have to estimate a priori error covariance matrix or prediction error covariance matrix that is given where P of n given $n-1$ and this is related to previous estimation error covariance matrix by this relationship.

This is the state matrix $A(n)$ into $P(n-1|n-1)$ into n transpose + Q_w , Q_w is the covariance matrix the process noise so we have the a priori error covariance matrix or prediction error covariance matrix. Now we will complete the Kalman gain that is given by this matrix K_n matrix that is related to the error covariance matrix P of n given $n-1$ multiplied by $c(n)$ transpose into inverse of this quantity $c(n) + p(n|n-1)$ into $c(n)$ transpose + Q_v whole inverse.

So, this is the expression for Kalman gain now once we have the Kalman gain we can now update this step that is $\hat{X}(n|n) = \hat{X}(n|n-1)$ which we have obtained here and then K_n times the prediction error $Y(n) - c(n)\hat{X}(n|n-1)$ this is the innovation or the prediction error. Now after this we have to update the a posteriori error covariance matrix that is given by $P(n|n) = I - K_n c(n)$ into $P(n|n-1)$.

This is the a priori error covariance matrix and this is pre multiplied by this vector to get the estimation error covariance matrix or a posteriorly error covariance matrix. So, this is the Kalman filter algorithm.

(Refer Slide Time: 31:16)



It is a simple algorithm, but it is widely used some applications are as follows Kalman filter has become an integral component of thousands of military and civilian navigation and guidance systems. So, any navigation or guidance system Kalman filter is a part. NASA adopted the Kalman filter to solve the problems associated with determining satellite orbits guided the Apollo 11 lunar module to the moon surface using this Kalman filter.

The Global Positioning System GPS is based on this deceptively simple algorithm. Similarly, it can be used for video based tracking time series based forecasting of financial data etc. There are many applications and there are many improved version of this algorithm.

(Refer Slide Time: 32:21)

Example

- ❖ Suppose a particle is moving with a random acceleration $a(n)$ which can be modelled as a white noise with variance 1. We have to estimate the position of the particle in presence of the measurement noise which is also white with variance 1
- ❖ Let $X_1(t)$ and $X_2(t)$ be the position and the velocity of the particle at time t . Assuming the sampling interval to be one second, we can approximately write

$$X_1(n) = X_1(n-1) + X_2(n-1)$$

$$X_2(n) = X_2(n-1) + a(n)$$
- ❖ The measured position of the particle is given by

$$Y(n) = X_1(n) + V(n)$$
 where $V(n)$ is the measurement noise..

Let us consider one example suppose a particle is moving with a random acceleration $a(n)$ which can be modeled as a white noise with variance 1. Acceleration is random with variance 1 we have to estimate the position of the particle in presence of the measurement noise which is also white with variance 1. Let us solve this problem let $X_1(t)$ and $X_2(t)$ be the position and the velocity of the particle at time t .

These two will be the state variables assuming this sampling interval to be one second, we can approximately write that is position $X_1(n)$ at time $n = X_1(n-1)$ that is the position at in time $n-1$ plus distance traveled in one second that is $X_2(n-1)$ that is the velocity multiplied by 1. Therefore, $X_1(n) = X_1(n-1) + X_2(n-1)$ and velocity at instant n and $X_2(n) = X_2(n-1)$ velocity at instant $n-1$ plus the acceleration into 1.

Because after one second, we are considering so that way this part we can write as $X_2(n) = X_2(n-1) + a(n)$. The measured position of the particle is given by $Y(n) = X_1(n) + V(n)$ this is the observed variable that is observed data $Y(n) = X_1(n) + V(n)$ where $V(n)$ is a measurement noise which is given to be white with variance 1.

(Refer Slide Time: 34:30)

Example

- We can write the state equation as

$$\begin{bmatrix} X_1(n) \\ X_2(n) \end{bmatrix} = \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} X_1(n-1) \\ X_2(n-1) \end{bmatrix} + \begin{bmatrix} 0 \\ a(n) \end{bmatrix}$$

- The observation equation is given by

$$Y(n) = X_1(n) + V(n) = \begin{bmatrix} 1 & 0 \end{bmatrix} \begin{bmatrix} X_1(n) \\ X_2(n) \end{bmatrix} + V(n)$$

$$\begin{aligned} Q_w &= E \left[\frac{W(n)}{W(n)} \frac{W(n)}{W(n)} \right] \\ &= E \left[\begin{bmatrix} 0 & a(n) \\ a(n) & 0 \end{bmatrix} \right] \\ &= E \left[\begin{bmatrix} 0 & 0 \\ 0 & a^2(n) \end{bmatrix} \right] = \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} \\ Q_v &= 1 \end{aligned}$$

So, we can write from these expressions $X_{1n} = X_{1n-1} + X_{2n-1}$ and $X_{2n} = X_{2n-1} + a_n$ we can write the matrix states equations X_1 and X_{2n} these are the states $= \begin{bmatrix} 1 & 1 \\ 0 & 1 \end{bmatrix}$ that is the system matrix multiplied by X_{1n-1} X_{2n-1} these are the previous states plus first equation we do not have any noise term. So, that is 0 and second equation we have a noise term a_n so this is the term corresponding to W_n .

So, this is the state equation the observation equation is given by $Y_n = X_{1n} + V_n$ because position plus noise and this we can write as because it is X_{1n} . So, we can write it in terms of a matrix that is there $\begin{bmatrix} 1 & 0 \end{bmatrix}$ matrix this is a C matrix multiplied by X_{1n} X_{2n} . So, these are the states if I multiply it $\begin{bmatrix} 1 & 0 \end{bmatrix}$ that is the C matrix into state vector X_{1n} X_{2n} then 1 into X_{1n} + 0 into X_{2n} and we will get X_{1n} and + V_n .

So, V_n is the noise and it will have a variance because only the single dimensional quantities it will have variance which is σ_v^2 which is equal to 1. Thus here we can write that Q_w that is a matrix because it is a of this if this is w and vector W_n into W_n transpose. So, that will be equal to E of 0 an transpose will be 0 an and we will get E of 0 and this will be 0 and this one will be 0 and here a square and if we take the expected value, we will get 1 here because variance is 1.

So, that way it is 0 0 0 1 so this is the error covariance matrix here, so we know Q_w we know this is my A which is constant not time varying this is my C matrix and you Q_v is now it one-dimensional quantity that is equal to σ_v^2 actually that is equal to 1. So, these are the parameters given with this we can estimate this state.

(Refer Slide Time: 37:54)

The image shows handwritten mathematical derivations for the Kalman filter initialization at $n=0$. The steps are as follows:

$$\hat{x}(-1) = 0 \quad P(-1) = \sigma^2 I = \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}$$

$$\begin{aligned} n=0 \quad P(0| -1) &= A P(-1) A' + Q_w \\ &= \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} + \begin{bmatrix} 0 & 0 \\ 0 & 1 \end{bmatrix} \\ &= \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \end{aligned}$$

$$\begin{aligned} K_0 &= P(0| -1) C' (C P(0| -1) C' + Q_v)^{-1} \\ &= \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} \left(\begin{bmatrix} 1 & 0 \end{bmatrix} \begin{bmatrix} 1 & 0 \\ 0 & 2 \end{bmatrix} \begin{bmatrix} 1 \\ 0 \end{bmatrix} + 1 \right)^{-1} \\ &= \begin{bmatrix} \frac{1}{2} \\ 0 \end{bmatrix} \end{aligned}$$

$$\begin{aligned} \hat{x}(0) &= \hat{x}(-1) + K_0 (y(0) - C \hat{x}(-1)) = \begin{bmatrix} \frac{2}{3} \\ \frac{1}{3} \end{bmatrix} \\ P(0|0) &= (I - K_0 C) P(0| -1) = \begin{bmatrix} \frac{1}{2} & 0 \\ 0 & 2 \end{bmatrix} \end{aligned}$$

So, initial value we can assume that $\hat{x}(-1)$ given -1 that is initially this = 0 and then P of -1 that is error covariance matrix given -1 that will be also equal to some positive matrix we have to consider so that where $\sigma^2 I$ where σ^2 is a positive quantity. We may take this is σ^2 make we may take one so we can take this as the identity matrix itself. So, this is 1 0 0 1 now starting with $n = 0$ we will first estimate the prediction error covariance matrix so that way we will estimate P of 0 given -1 .

So, this is given by A into P - 1 given -1 this is the initial estimate of the error covariance matrix into A transpose + Q_w this is the expression. So, this will obtain from this expression A_n into P_{n-1} given $n-1$ into A_n transpose + Q_w . So, we get this expression and we can carry out the computation now this = 1 1 0 1 this is my A and then this is identity matrix 1 0 0 1 and this is A transpose that way it will be 1 0 1 1 + Q_w , Q_w matrix we have found out Q_w matrix is 0 0 0 1.

So, that way it will be 0 0 0 0 0 0 1 so we can carry out this computation and we will get this as 2 1 1 1 so this is the a priori error covariance matrix. Now when this we can find out K_0 Kalman

gain at time 0 = $P_{0|0} - 1 \text{ into } c^T \text{ into } c \text{ into } P_{0|0} - 1 \text{ into } c^T + Qv$ whole inverse. So, this is the expression we get from this formula, so this is the Kalman gain. So, using this formula we get this, so we have already obtained this, and we have value of c , c is $1 \ 0$ here c is known this one is known here c^T we know Qv is simply 1.

So, that way we can write this as this is this matrix $2 \ 1 \ 1 \ 1$ into c^T is c is $1 \ 0$. So, this will be $1 \ 0$ and here similarly this will be $1 \ 0$ this is the matrix then this is $2 \ 1 \ 1 \ 1$ and then c^T will be $1 \ 0 + Qv = 1$ and whole inverse so this we can carry out and we will get the here this matrix this is a 2×1 matrix two third and one third. Now we can estimate $x_{0|0}$ given 0 that is will be equal to $x_{0|0} - 1$. So plus Kalman gain that is k_0 matrix into $y_0 - c^T x_{0|0}$ times $x_{0|0}$ given that is y_0 here $x_{0|0} - 1$.

So, that way we can write, and this part is 0 because there is no prediction initially this part is also equal to 0 and therefore, we will simply have k_0 , k_0 is two third one third into y_0 . So, that way this will be two third of y_0 in and one-third of y_0 so this is the estimator for the states this is the position this is the velocity and now we can calculate $P_{0|0}$ that will be equal to $I - I$ had to be identity matrix into k_0 Kalman gain matrix into c all are given into $P_{0|0} - 1$.

So, we know all this quantity, so this is the identity matrix $1 \ 0 \ 0 \ 1$ and then k_0 we know that is equal to two third one third this matrix then c we know c is $1 \ 0$ and then this $P_{0|0} - 1$ that also we know that is given by $2 \ 1 \ 1 \ 1$. So, if we carry out this calculation, we will get this as the two third, one third, one third, two third so that way we can find out $P_{0|0}$ then we can find out $P_{1|1}$ given 0 then $k_1 \hat{x}_{1|1}$ given 1 like that. So, if we can estimate the state in this recursive manner.

(Refer Slide Time: 45:21)

Summary

❖ The vector Kalman filter is a recursive state estimator, for the state space model

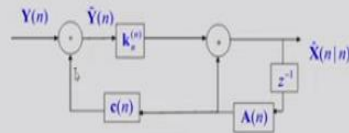
$$\mathbf{X}(n) = \mathbf{A}(n)\mathbf{X}(n-1) + \mathbf{W}(n) \quad (\text{state equation})$$

$$\mathbf{Y}(n) = \mathbf{c}(n)\mathbf{X}(n) + \mathbf{V}(n) \quad (\text{observation equation})$$

❖ $\mathbf{W}(n)$ is assumed to be a white noise vector. Thus, the covariance matrix $\mathbf{Q}_w$ is a diagonal matrix

❖ $\mathbf{V}(n)$ is also assumed to be a white noise vector with the covariance $\mathbf{Q}_v$. Further, $\mathbf{V}(n)$ is independent of $\mathbf{X}(n)$ and $\mathbf{W}(n)$.

❖ The Kalman filter structure is given by the following block diagram



Let us summarize the lecture the vector Kalman filter is a recursive state estimator for the state space model this is this state space model $\mathbf{X}_n = \mathbf{A}_n \mathbf{X}_{n-1} + \mathbf{W}_n$ $\mathbf{A}_n$ is the state matrix or system matrix $\mathbf{X}_n$ is the state vector $\mathbf{W}_n$ is the process noise vector $\mathbf{Y}_n$ that is the measurement vector or output data vector $= \mathbf{c}_n \mathbf{X}_n + \mathbf{V}_n$ that is known as the output matrix multiplied by the state vector $+ \mathbf{V}_n$, $\mathbf{V}_n$ is the measurement noise vector.

$\mathbf{W}_n$ is assumed to be white noise vector thus the covariance matrix $\mathbf{Q}_w$ is a diagonal matrix $\mathbf{V}_n$ is also assumed to be a white noise vector with the covariance $\mathbf{Q}_v$. Further $\mathbf{V}_n$ is independent of $\mathbf{X}_n$ and $\mathbf{W}_n$ that is important the Kalman filter structure is given by the following block diagram. So, this is the input data $\mathbf{Y}_n$ and this is the $\hat{\mathbf{X}}_n$ given n this is the output state estimator. So, this is the input that is the data vector.

Now this is delayed that will give us $\hat{\mathbf{X}}_{n-1}$ given $n-1$ that will be multiplied by $\mathbf{A}_n$ and we will get the predicted part here and this predicted part is multiplied by $\mathbf{c}_n$ to get the predicted data here we get the predicted state here, we get the predicted data then this is subtracted we get the innovation. This innovation we are scaling by the Kalman matrix and this is added to the predicted part to get the state estimator.

(Refer Slide Time: 47:30)

Summary...

- ❖ Following the same procedure as the scalar Kalman filter, we can derive the following Kalman filter equations:

- (1) Obtain Update the prediction error covariance matrix

$$\mathbf{P}(n|n-1) = \mathbf{A}(n)\mathbf{P}(n-1|n-1)\mathbf{A}'(n) + \mathbf{Q}_w$$

- (2) Estimate Kalman gain

$$\mathbf{k}_n = \mathbf{P}(n|n-1)\mathbf{c}'(n)(\mathbf{c}(n)\mathbf{P}(n|n-1)\mathbf{c}'(n) + \mathbf{Q}_v)^{-1}$$

Estimate states

$$\hat{\mathbf{X}}(n|n) = \mathbf{A}(n)\hat{\mathbf{X}}(n-1|n-1) + \mathbf{k}_n(Y(n) - \mathbf{c}(n)\mathbf{A}(n)\hat{\mathbf{X}}(n-1|n-1))$$

- (3) Update the estimation error covariance matrix

$$\mathbf{P}(n|n) = (\mathbf{I} - \mathbf{k}_n\mathbf{c}(n))\mathbf{P}(n|n-1)$$

- ❖ Kalman filter has diverse applications

Following this same procedure at this scalar Kalman filter we can derive the following Kalman filter equations that is obtain the update of the prediction error or a priori error covariance matrix $\mathbf{P}$ of n given $n-1 = \mathbf{A}(n) \text{ into } \mathbf{P} \text{ of } n-1 \text{ given } n-1 \text{ into } n \text{ transpose} + \mathbf{Q}_w$. Estimate the Kalman gain that is $\mathbf{K}_n = \mathbf{P} \text{ of } n \text{ given } n-1 \text{ into } \mathbf{c}(n) \text{ transpose into inverse of this that is } \mathbf{c}(n) \text{ into } \mathbf{P} \text{ of } n \text{ given } n-1 \text{ into } \mathbf{c}(n) \text{ transpose} + \mathbf{Q}_v$ where $\mathbf{Q}_v$ is the observation noise.

This is the covariance matrix of the observation noise then we will estimate these states that is $\hat{\mathbf{X}}(n|n) = \mathbf{A}(n) \text{ into } \hat{\mathbf{X}}(n-1|n-1)$ this is the prediction part + $\mathbf{K}_n$ Kalman gain into one correction part that is $Y(n) - \mathbf{c}(n) \text{ into } \mathbf{A}(n) \text{ into } \hat{\mathbf{X}}(n-1|n-1)$. This is the previous estimation now we will update the estimation error because here estimation error is in use. So, for that we have to update it and this update is given by $\mathbf{I} - \mathbf{K}_n \text{ into } \mathbf{c}(n) \text{ into } \mathbf{P} \text{ of } n \text{ given } n-1$ this is the covariance matrix update equation. So, this is the Kalman filter algorithm and we are also stated that Kalman filter has diverse applications. Thank you.