

Speech coding

①

- Requirement of speech coding :
 - a) high speech quality
 - b) low bit rate
 - c) low complexity
 - d) limited signal delay.

Note: some contradicting requirements. we need trade-off.

- Speech coding algorithms can be mainly divided into three categories:
 - a) waveform coding
 - b) parametric coding
 - c) hybrid coding

Waveform coding : quantization of samples of speech signal itself or some intermediate signals, such as prediction error signal or subband signals.

Example: DPCM, ADPCM

quality: good, SNR = 30-35 dB

bit rate: $f_s = 8 \text{ kHz}$
 $B = 32 \text{ kbits/sec}$

parametric coders (vocoders): Do not code the waveform, but a set of model parameters. The model parameters are quantized and transmitted. Using the coded model parameters, speech is generated at the receiver.

Example: Vocoder

Application area: mainly for defence, military deployed in a no-man's land.

quality: synthetic speech

bit rate: $B = 2.4 \text{ Kbits/sec}$

②

hybrid coding: ① A compromise between waveform coding and parametric coding.

② The parameters of a time-varying speech signal is quantized and send. Also the residual signal is quantized blockwise and send.

Example: ~~many~~ ^{many} telephone quality speech codes.

bitrate: 0.75 - 1.5 bits/sample.

quality: SNR ~ 10 dB

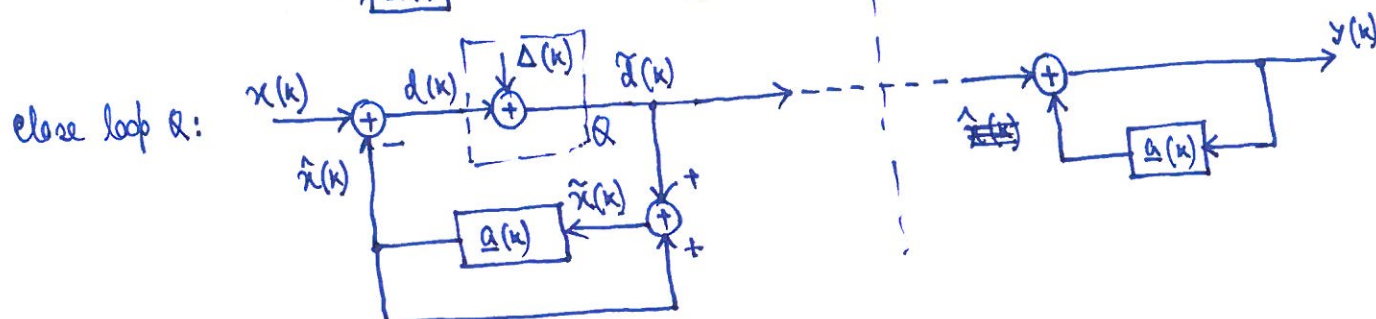
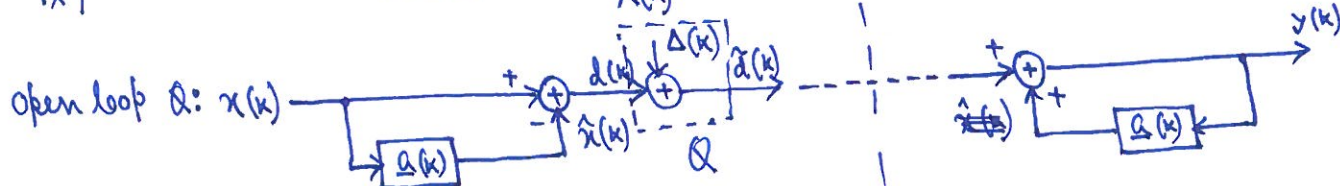
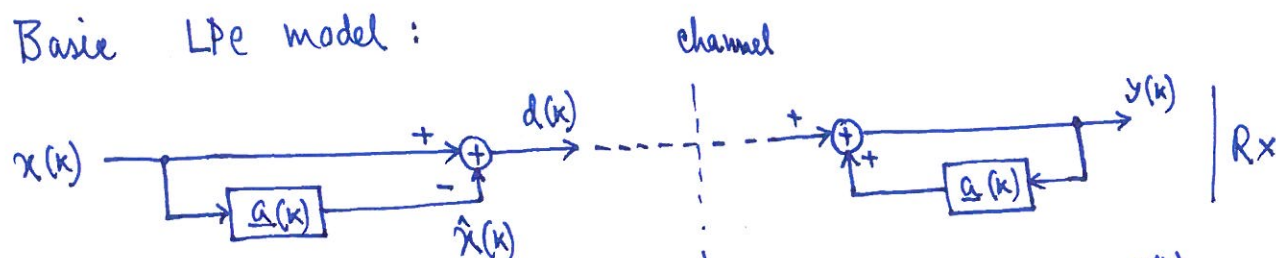
subjective quality: OK

Note: All the three approach uses linear Prediction (LP).

Generic coding term: Linear Predictive Coding (LPC).

Quantization of the Residual Signal

Basic LPE model:



Quantization with open-loop prediction

Quantized
residual:

$$\tilde{d}(k) = d(k) + \Delta(k).$$

$$\tilde{D}(z) = D(z) + \Delta(z).$$

Now, $\tilde{D}(z) + A(z)Y(z) = Y(z)$

So, $\frac{Y(z)}{\tilde{D}(z)} = \frac{1}{1-A(z)}.$

or, $Y(z) = \frac{1}{1-A(z)} \tilde{D}(z) = \frac{1}{1-A(z)} D(z) + \frac{1}{1-A(z)} \Delta(z).$

What is $D(z)$ and $X(z)$ relation?

$$\begin{aligned} D(z) &= X(z) - A(z)X(z) \\ &= (1-A(z))X(z) \end{aligned}$$

or, $\frac{D(z)}{1-A(z)} = X(z)$

$\therefore Y(z) = X(z) + \frac{1}{1-A(z)} \Delta(z).$

or, $y(k) = x(k) + r(k)$ where $r(k) = \mathcal{Z}^{-1} \left(\frac{1}{1-A(z)} \Delta(z) \right).$

Assume that the ~~noise~~ quantization noise $\Delta(k)$ is ~~uniformly~~ ^{white} uniformly distributed as $-\frac{\Delta d}{2} \leq \Delta(k) \leq +\frac{\Delta d}{2}$ (in range)

with the power $P_{\Delta} = \frac{(\Delta d)^2}{12}$

and a constant noise PSD

$$\phi_{\Delta\Delta}(e^{j\omega}) = \frac{(\Delta d)^2}{12}. \quad \dots \quad (i)$$

Here, Δd denotes the quantizer's step size.

Note : ① The reconstruction error is $r(k)$.

② The $r(k)$ is a filtered version of $\Delta(k)$, the quantization noise.
 ~~$\Delta(k)$ is white~~

③ The spectrum of the ~~quantization noise~~ reconstruction error $r(k)$ follows the spectral envelope of speech signal.

Why? Because $\frac{1}{1-A(z)}$ is the synthesis filter that models the vocal tract.

④ This is called "noise shaping".

- Let us see the interplay between prediction gain, signal-to-reconstruction-error ratio, signal-to-quantization-noise ratio.

Signal-to-reconstruction error ratio = Signal-to-noise ratio for output signal. $\left(\frac{S}{N}\right)_y$.

$$\left(\frac{S}{N}\right)_y = \frac{\sigma_x^2}{\sigma_r^2} = \frac{\phi_{xx}(0)}{\phi_{rr}(0)} \quad \cdot \text{ we don't know } \phi_{rr}(0).$$

Assumption : ① Let us assume the LP provides a perfect job.

② The LP prediction error $d(k)$ is white.

$$\textcircled{c} \text{psd: } \phi_{dd}(e^{j\omega}) = \text{const} = \phi_{dd}(0) = \sigma_d^2 \dots \textcircled{ii}$$

$$\text{Now, } \phi_{dd}(e^{j\omega}) = \phi_{xx}(e^{j\omega}) \cdot |1 - A(e^{j\omega})|^2$$

$$\text{w, } \frac{\phi_{dd}(0)}{|1 - A(e^{j0})|^2} = \phi_{xx}(e^{j0}) \quad [\text{using ii}]$$

$$\text{w, } \phi_{dd}(0) \cdot \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1}{|1 - A(e^{j\omega})|^2} d\omega = \frac{1}{2\pi} \int_{-\pi}^{\pi} \phi_{xx}(e^{j\omega}) d\omega = \phi_{xx}(0) = \sigma_x^2.$$

(5)

$$\therefore \frac{\sigma_n^2}{\sigma_d^2} = \frac{\varphi_{nn}(0)}{\varphi_{dd}(0)} = \frac{1}{2\pi} \int_{-\pi}^{\pi} \frac{1}{|1 - A(e^{j\omega})|^2} d\omega = G_p \dots \text{iii)}$$

$$\text{Define } |H(e^{j\omega})|^2 = \frac{1}{|1 - A(e^{j\omega})|^2}$$

Now, Signal-to-quantization-noise-ratio = Signal-to-noise ratio for prediction error $d(k)$

$$\left(\frac{S}{N}\right)_x = \frac{\sigma_d^2}{\sigma_\Delta^2} = \frac{\varphi_{dd}(0)}{\left(\frac{(\Delta d)^2}{12}\right)} \dots \text{(iv)}$$

$$\begin{aligned} \text{Now, } \phi_{rr}(e^{j\omega}) &= |H(e^{j\omega})|^2 \cdot \phi_{\Delta\Delta}(e^{j\omega}) \\ &\quad \left(\text{because } \sum \{r(k)\} = \frac{1}{1-A(z)} \Delta(z) \right) \\ R(z) &= \frac{1}{1-A(z)} \cdot \Delta(z) \end{aligned}$$

$$\begin{aligned} \therefore \varphi_{rr}(0) &= \frac{1}{2\pi} \int_{-\pi}^{\pi} \phi_{rr}(e^{j\omega}) d\omega \\ &= \frac{1}{2\pi} \int_{-\pi}^{\pi} |H(e^{j\omega})|^2 \phi_{\Delta\Delta}(e^{j\omega}) d\omega \\ &= \frac{(\Delta d)^2}{12} \cdot \frac{1}{2\pi} \int_{-\pi}^{\pi} |H(e^{j\omega})|^2 d\omega \quad [\text{using (ii)}] \\ &= \frac{(\Delta d)^2}{12} \cdot G_p \quad [\text{using (iii)}] \end{aligned}$$

$$\therefore \left(\frac{S}{N}\right)_y = \frac{\sigma_n^2}{\sigma_r^2} = \frac{\varphi_{nn}(0)}{\varphi_{rr}(0)} = \frac{G_p \cdot \varphi_{dd}(0)}{\frac{(\Delta d)^2}{12} \cdot G_p} = \frac{\varphi_{dd}(0)}{\frac{(\Delta d)^2}{12}} = \left(\frac{S}{N}\right)_x$$

(using (iii) and (iv)).

Main point: The ultimate quality is the $\left(\frac{S}{N}\right)_y$, which did not improve by open-loop prediction. But we get spectral shaping.

Closed loop Quantization

⑥

$$\begin{aligned} y(k) &= \hat{x}(k) = \hat{x}(k) + \tilde{d}(k) \\ &= \hat{x}(k) + d(k) + \Delta(k) \\ &= x(k) + \Delta(k) . \\ &\quad \uparrow \\ &\quad d(k) = x(k) - \hat{x}(k) \end{aligned}$$

Note : (a) The reconstruction error is the quantization noise $\Delta(k)$.

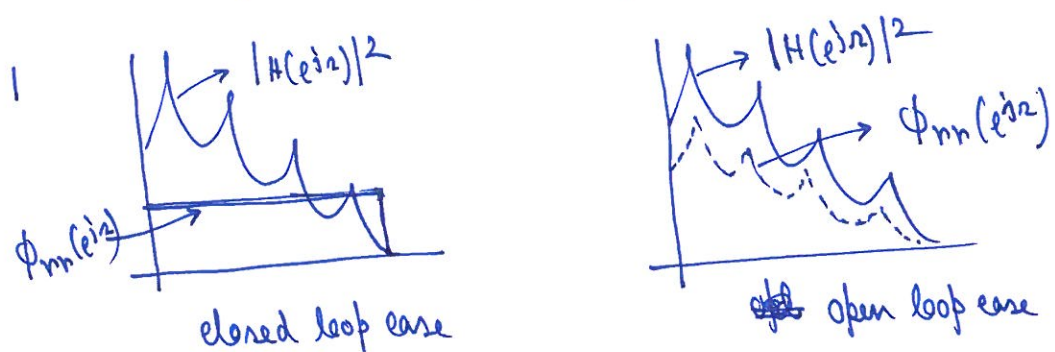
(b) The reconstruction error is spectrally white.

$$\begin{aligned} \left(\frac{S}{N}\right)_y &= \frac{\sigma_x^2}{\sigma_\Delta^2} = \frac{\sigma_x^2}{\sigma_d^2} \cdot \frac{\sigma_d^2}{\sigma_\Delta^2} \\ &= G_p \cdot \left(\frac{S}{N}\right)_d \end{aligned}$$

$$\therefore \text{SNR}_y = 10 \log_{10} G_p + 10 \log_{10} \left(\frac{S}{N}\right)_d$$

- Closed loop prediction (following 6-dB-per-bit rule), provides $\frac{10 \log_{10} G_p}{6}$ bit advantage than direct uniform quantization of speech signal (PEM).
- But, we did not get any spectral shaping.

- See Figure 8.10 of the reference book



Main points:	<u>closed loop</u>	<u>Open loop</u>
spectral shaping	X	✓
SNR improvement	✓	X

But, we want both.

Q: Can we do in between? A scheme that combined closed loop and open loop.

Ans: Yes. It is possible by generalizing the closed loop structure and inserting a noise ~~shaping~~ filter.