

Modelling the Arrival Process for Packet Audio

Ingemar Kaj
Dept. of Mathematics
Uppsala University
Sweden
ikaj@math.uu.se

Ian Marsh
CNA Lab.
SICS
Sweden
ianm@sics.se

Abstract

Packet audio streams can be distorted during the traversal of a packet switched network. The inter-packet spacing of the sender stream is not preserved as packets encounter queues in routers. The contribution of this work is a Markov model for the time delay variation of packetised voice traffic. Our model is extensible and we show this by adding silence suppression as well as including packet loss. By comparing the model to real traffic traces taken, we show it is possible to generate a packet audio arrival process very similar to real conditions. We illustrate this by comparing the probability density functions for our model and the captured trace data.

Keywords:: Packet delay, VoIP, Markov chain, Steady state

1 Introduction

Modelling the arrival process of audio packets that have passed through between a number of routers is the problem we will address. Figure 1 illustrates the scenario, packets containing audio samples are sent at a constant rate from a sender in stage 1.

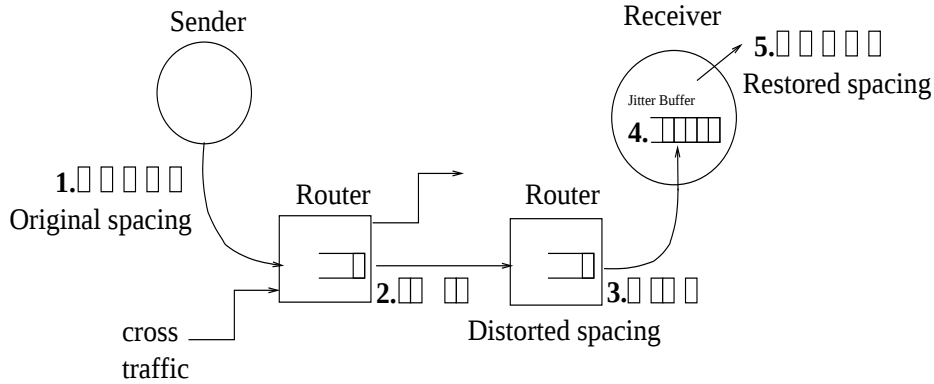


Figure 1: Spacing distortion due to queueing

The packets are compressed and elongated relative to each other, mainly due to the buffering in intermediate routers and mixing with cross-traffic, shown in stages 2 and 3. Finding mathematical parameters that model the compression and elongation of packets is the problem we will address in this paper.

The standard method to solve this problem is to introduce a small buffer in the end systems between the decoded audio stream and the hardware. This is shown in the figure and is often called a jitter buffer, stage 4. Some additional timing information is needed in each packet indicating to the application when the particular packet should be replayed. Using this information the original packet spacing can be restored, shown in stage 5. A protocol exists for including this timing information and is known as the Real Time Protocol (RTP).

The motivation for this work derives from the inability of using known arrival processes to approximate the packet audio arrival process at the receiver, and the inherent static nature of using pre-recorded trace files. Using a known arrival process, even a complex one, is not always realistic as the model does not include characteristics that real audio data experiences, for example the use of silence suppression or the contribution of cross traffic. Alternatively, using traces

taken from sessions produces accurate and representative arrival processes, but not the flexibility needed when testing jitter buffer performance for example. The traces can only be used statically, for example, it is not possible to vary the sending frequency. Therefore an important contribution of this paper is to address the deficiencies of these two approaches to create a more representative model of packet audio arrivals by *combining* the advantages of the two approaches. Our long term goal is to use the model for stressing jitter buffer playout algorithms.

A secondary goal is to obtain a better insight into the arrival process of packet audio streams. The effect of queueing systems on constant rate traffic has been studied in some detail but with the goal of calculating buffer requirements or obtaining admission control thresholds. Modelling the arrival of a single stream through a sequential system of buffers has not been explored with our stated goals.

This paper presents in a descriptive manner a packet delay model, based on the main assumption that packets are subject to independent network delays. It is intended that readers not completely familiar with Markovian theory can still follow the description using the text and graphs as well as the mathematical details. Little prior knowledge of the system is assumed, the model is built from first principles in Section 2 where the model is introduced. We concentrate first on mean values and continue with the steady state behaviour. We add silence suppression in Section 3, packet loss in Section 4, incorporate real sample data in Section 5, and give an overview of the results in Section 5.2. We present some related efforts in Section 6 and finally round off with some conclusions.

2 Packet delay model

Consider an audio stream sent from a single source where equal size voice packets are transmitted periodically with a 20 ms packetisation time. For convenience the packetisation interval is taken to be the time unit in the model. Saying that a packet is sent at time t means that the packet is sent at clock time $20t$ ms into the data stream. The first packet is sent at time 0.

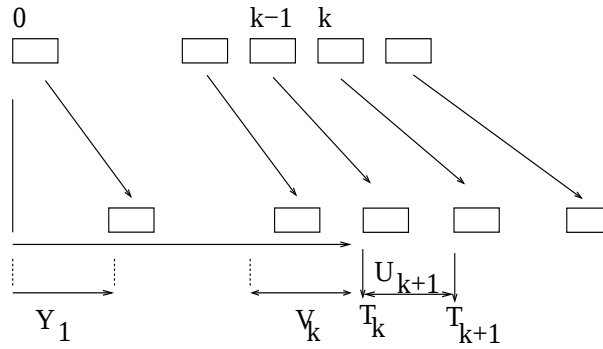


Figure 2: Model Representation

We begin with the network delay for a packet. Suppose that packet k could be sent isolated from the rest of the audio stream and let

$$Y_k = \text{network delay of packet } k \text{ (no. of 20 ms periods)}.$$

The network delay signifies the effect caused by interaction with cross traffic and the effect of buffers en-route to the receiving destination. As an approximation we assume that the network delays are described by a sequence of positive i.i.d. random variables $Y_1, Y_2, \dots$ with distribution function

$$F(x) = P(Y_k \leq x), \quad k \geq 1,$$

and mean network delay $\nu = \int_0^\infty (1 - F(x)) dx < \infty$. For the data in our study, typical values of ν are 20-40, i.e. 400-800 ms.

Now let

$$T_k = \text{the arrival time at receiver of packet } k, \quad k \geq 1.$$

The packet spends the time Y_k propagating through the network and arrives at the receiver (stage 3 in Figure 1) at time $k - 1 + Y_k$, unless it catches up one or several of the previously transmitted packets numbered 1 to $k - 1$ and in this way is forced to adapt to the slower pace of packets ahead in line. Since packets usually arrive in the same order as they are sent, these considerations show that the actual arrival times form a (transient) Markov chain, such that

$$T_1 = Y_1, \quad T_k = \max(T_{k-1}, k - 1 + Y_k), \quad k \geq 2. \quad (1)$$

It follows directly that the observed packet interarrival times are obtained from

$$U_k = T_k - T_{k-1} = \max(0, k - 1 + Y_k - T_{k-1}) \quad k \geq 2. \quad (2)$$

We introduce in addition the *observed delay* of the packets as

$$V_k = \text{arrival time} - \text{departure time} = T_k - k + 1 \quad k \geq 1.$$

In Section 5.1 below we study these quantities for measured data traces. As an example, Figure 3 shows the histogram for

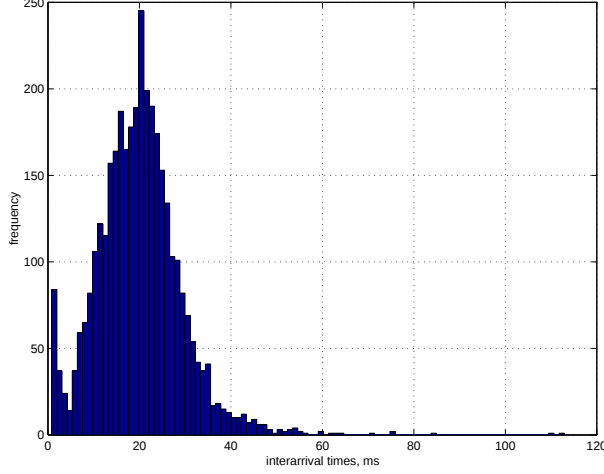


Figure 3: Histogram of interarrival times

an empirical sequence of interarrival times, corresponding to the sequence (U_k) . The data is from a recording of a Voice over IP (VoIP) session between Argentina and Sweden.

2.1 Mean arrival and interarrival times

It is intuitively clear that in the long run $E(U_k) \approx 1$ since on average packets arrive with 20 ms spacings, which we now verify for the model. The representation (1) for T_k can be written

$$T_k = \max(Y_1, 1 + Y_2, \dots, k - 1 + Y_k) \quad k \geq 1,$$

which gives the alternative representation

$$T_k = \max(Y_1, 1 + T'_{k-1}), \quad k \geq 2 \quad (3)$$

where on the right side

$$T'_{k-1} = \max(Y_2, 1 + Y_3, \dots, k - 2 + Y_k)$$

has the same marginal distribution as T_{k-1} but is independent of Y_1 . From (3) follows that we can write $\{T_k > t\}$ as a union of two disjoint events, as

$$\{T_k > t\} = \{1 + T'_{k-1} > t\} \cup \{Y_1 > t, 1 + T'_{k-1} \leq t\}.$$

Hence, using the independence of T'_{k-1} and Y_1 ,

$$\begin{aligned} P(T_k > t) &= P(1 + T'_{k-1} > t) + P(Y_1 > t, 1 + T'_{k-1} \leq t) \\ &= P(1 + T_{k-1} > t) + P(Y_1 > t)P(1 + T_{k-1} \leq t) \end{aligned}$$

and so

$$\begin{aligned} E(T_k) &= \int_0^\infty P(T_k > t) dt \\ &= E(1 + T_{k-1}) + \int_1^\infty P(Y_1 > t)P(T_{k-1} \leq t - 1) dt. \end{aligned} \quad (4)$$

Therefore

$$E(U_k) = 1 + \int_1^\infty P(Y_1 > t)P(T_{k-1} \leq t-1) dt \rightarrow 1, \quad k \rightarrow \infty \quad (5)$$

(since $\nu = \int_0^\infty P(Y_1 > t) dt < \infty$ and $T_k \rightarrow \infty$, the dominated convergence theorem applies forcing the integral to vanish in the limit).

A further consequence of (4) is obtained by iteration,

$$E(T_k) = k - 1 + E(Y_1) + \int_1^\infty P(Y_1 > t) \sum_{i=1}^{k-1} P(T_i \leq t-1) dt.$$

If we introduce

$$N(t) = \text{the number of arriving packets in the time interval } (0, t],$$

so that $\{N(t) \geq n\} = \{T_n \leq t\}$, this can be written

$$E(V_k) = E(Y_1) + \int_1^\infty P(Y_1 > t) \sum_{i=1}^{k-1} P(N(t-1) \geq i) dt, \quad (6)$$

which, as $k \rightarrow \infty$, gives an asymptotic representation for the average observed delay as

$$E(V_k) \rightarrow \nu + \int_1^\infty P(Y_1 > t)E(N(t-1)) dt. \quad (7)$$

2.2 Steady state distributions

By (1),

$$P(T_k \leq x) = \prod_{i=1}^k P(i + Y_i \leq x + 1) = \prod_{i=0}^{k-1} F(x - i),$$

and therefore the sequence (V_k) , which we defined by $V_k = T_k - k + 1$, $k \geq 1$, satisfies

$$P(V_k \leq x) = \prod_{i=0}^{k-1} F(x + k - 1 - i) = \prod_{i=0}^{k-1} F(x + i) \quad x \geq 0.$$

This shows that (V_k) is a Markov chain with state space the positive real line and asymptotic distribution given by

$$P(V_\infty \leq x) = \prod_{i=0}^\infty F(x + i) \quad x \geq 0. \quad (8)$$

Furthermore, for $x \geq 0$

$$\begin{aligned} P(U_k \geq x) &= P(k - 1 + Y_k - T_{k-1} \geq x) = P(V_{k-1} \leq Y_k + 1 - x) \\ &= \int_0^\infty P(V_{k-1} \leq y + 1 - x) dF(y), \end{aligned}$$

where in the step of conditioning over Y_k we use the independence of Y_k and V_{k-1} . Therefore the sequence (U_k) has the asymptotic distribution

$$P(U_\infty \leq x) = 1 - \int_0^\infty \prod_{i=1}^\infty F(y - x + i) dF(y) \quad x \geq 0, \quad (9)$$

in particular a point mass in zero of size

$$P(U_\infty = 0) = 1 - \int_0^\infty \prod_{i=1}^\infty F(y + i) dF(y). \quad (10)$$

This distribution has the property that $E(U_\infty) = 1$ for any given distribution F of Y with $\nu = E(Y) < \infty$. In fact, this follows from 5 under a slightly stronger assumption on Y (uniform integrability) but can also be verified directly by integrating (9). Figure 4 shows numeric approximations of the (non-normalised) density function $\frac{d}{dx}P(U_\infty \leq x)$ of (9) for three choices of F , one Gaussian and two exponential distributions. A fraction of the probability mass is fixed at $x = 0$ in accordance with (10).

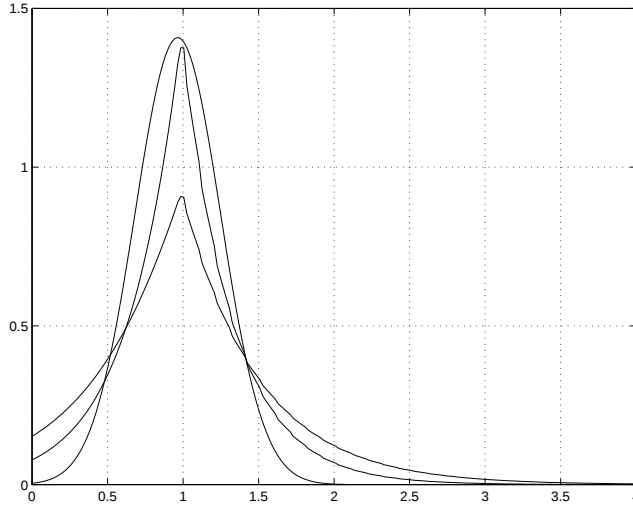


Figure 4: Density functions of U for $N(5,0.2)$, $\text{Exp}(2)$ and $\text{Exp}(3)$

3 Silence suppression mechanism

In this section we incorporate the additional source of random delays due to silence suppression into the model. Silence suppression is used to save sending when the talker is silent. During a normal conversation, this accounts for about half of the total number of packets, considerably reducing the load on a network. Assign to packet no k the quantity

$$X_k = \text{duration of silent period between packets no } k-1 \text{ and } k.$$

A silent period is the time interval during which the silence suppression mechanism is in effect. We assume that the silence suppression intervals are independent of $(Y_k)_{k \geq 1}$ and given by a sequence of independent random variables $X_1, X_2, \dots$, such that

$$G(x) = P(X_k \leq x), \quad 1 - \alpha = G(0) = P(X_k = 0) > 0, \quad \mu = E(X_k) < \infty.$$

The (small) probability $\alpha = P(X_k > 0)$ represents the case where silence suppression is activated just after packet $k-1$ is delivered from the sender. Note that

$$S_k = \sum_{i=1}^k X_i = \text{total time of silence suppression affecting packet } k,$$

which implies that the delivery of packet k from the sending unit now starts at time $k-1 + S_k$. The representation (1) takes the form

$$T_1 = S_1 + Y_1, \quad T_k = \max(T_{k-1}, k-1 + S_k + Y_k), \quad k \geq 2, \quad (11)$$

hence

$$U_k = T_k - T_{k-1} = \max(0, k-1 + S_k + Y_k - T_{k-1}) \quad k \geq 2. \quad (12)$$

Similarly,

$$V_k = \text{arrival time} - \text{departure time} = T_k - k + 1 - S_k \quad k \geq 1.$$

The alternative representation (3) is

$$T_k = X_1 + \max(Y_1, 1 + T'_{k-1}), \quad (13)$$

where

$$T'_{k-1} = \max(Y_2 + S_2 - X_1, 1 + Y_2 + S_2 - X_1, \dots, k-2 + Y_k + S_k - X_1)$$

has the same marginal distribution as T_{k-1} but is independent of X_1 and Y_1 . In analogy with the calculation of the previous section leading up to (4), this relation gives

$$E(T_k) = E(X_1 + 1 + T_{k-1}) + \int_1^\infty P(X_1 + Y_1 > t, X_1 + T'_{k-1} \leq t-1) dt. \quad (14)$$

Exchanging the operations of integration and expectation shows that the last integral can be written

$$E \left[\int_{1+X_1}^{\infty} \mathbf{1}\{Y_1 > t - X_1, T'_{k-1} > t - X_1 - 1\} dt \right]$$

where we have also used that the integrand vanishes on the set $\{t \leq 1 + X_1\}$. Apply the change-of-variables $t \rightarrow t - X_1$ to get $E \left[\int_1^{\infty} \mathbf{1}\{Y_1 > t, T'_{k-1} > t - 1\} dt \right]$. Then shift integration and expectation again to obtain from (14) the relations

$$E(T_k) = 1 + E(X_1) + E(T_{k-1}) + \int_0^{\infty} P(Y_1 > t)P(T_{k-1} \leq t - 1) dt$$

and

$$E(U_k) = 1 + E(X_1) + \int_1^{\infty} P(Y_1 > t)P(T_{k-1} \leq t - 1) dt.$$

Hence with silence suppression, as $k \rightarrow \infty$,

$$E(U_k) \rightarrow 1 + \mu, \quad E(V_k) \rightarrow \nu + \int_1^{\infty} P(Y_1 > t)E(N(t - 1)) dt, \quad (15)$$

using the same arguments as in the simpler case of the previous section.

4 Including packet loss in the model

We return to the original model without silence suppression but consider instead the effect of lost packets. Suppose that each IP packet is subject to loss with probability p , independently of other packet losses and of the network delays. Lost packets are unaccounted for at the receiver and hence, in this section, the sequence (T_k) records arrival times of non-lost packets only. To keep track of their delivery times from the sender introduce

$$K_k = \begin{array}{l} \text{number of attempts required between} \\ \text{successful packets } k - 1 \text{ and } k, \quad k \geq 1, \end{array}$$

which gives a sequence $(K_k)_{k \geq 1}$ of independent, identically distributed random variables with the geometric distribution

$$P(K_k = j) = (1 - p)p^j, \quad j \geq 0.$$

Moreover,

$$\begin{aligned} L_k &= K_1 + \dots + K_k \\ &= \text{number of attempts required for } k \text{ successful packets} \end{aligned}$$

is a sequence of random variables with the negative binomial distribution. The arrival times of packets are now given by

$$T_1 = K_1 - 1 + Y_{K_1}, \quad T_k = \max(T_{k-1}, L_k - 1 + Y_{L_k}), \quad k \geq 2.$$

Due to the independence we may re-index the sequence of Y_{L_k} 's to obtain

$$T_1 = K_1 - 1 + Y_1, \quad T_k = \max(T_{k-1}, L_k - 1 + Y_k), \quad k \geq 2.$$

and thus

$$T_k = K_1 - 1 + \max(Y_1, 1 + T'_{k-1}), \quad k \geq 2 \quad (16)$$

with K_1 , Y_1 and T'_{k-1} all independent, and again T_{k-1} and T'_{k-1} identically distributed. This is the same relation as (13) with X_1 replaced by $K_1 - 1$ and hence, as in (15), $E(U_k) \rightarrow 1 + E(K_1 - 1) = \frac{1}{1-p}$, $k \rightarrow \infty$, which provides a simple method to estimate packet loss based on observed interarrival times. Similarly, combining silence suppression and packet loss,

$$E(U_k) \rightarrow 1 + E(X) + E(K_1 - 1) = \mu + \frac{1}{1-p}, \quad k \rightarrow \infty, \quad (17)$$

5 Incorporating Real Data

5.1 Trace data

We give a brief description of the experiments we performed in order to obtain estimates for the parameters in the model. Packet audio streams were sent from a site in Buenos Aires, Argentina to Stockholm, Sweden over a number of weeks. The packet streams were recorded as trace files¹ at the receiver. The remote sending site is approximately 12,000 kilometres (7500 miles), 25 Internet hops and four time zones from our receiver.

We use our own packet audio tool, Sicsophone, for sending a constant rate 64kbits per second, Pulse Code Modulation (PCM) stream of 160 byte audio packets. This implies the packets leave the sender with a inter-packet spacing of 20ms. An additional 12 bytes of RTP information is added to this data which we use for the detection of lost packets and the start and end of talkspurts. Sicsophone is capable of silence suppression, in which packets are not sent where no detectable audio is found. This avoids unnecessary capacity waste. A pre-recorded conversation is used as the content with silence periods. Without silence suppression, 3563 packets are sent during 70 seconds and with suppression 2064 are sent. We record the absolute times the packets leave the sender and the absolute arrival times at the receiver. This gives an observed

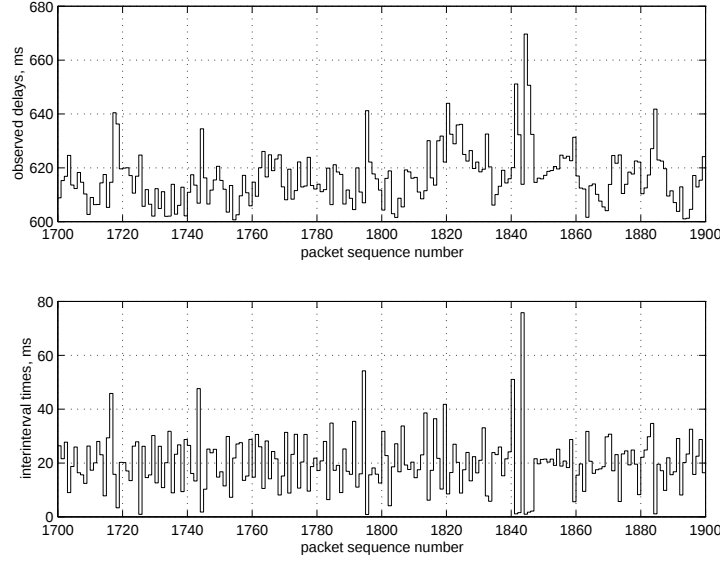


Figure 5: Four second audio packet trace; delay and interarrival times

sequence

$$v_k = \text{arrival time no } k - \text{departure time no } k$$

of the Markov chain (V_k) . In particular, the sample mean $\bar{v}$ is an estimate of the one-way delay. Similarly,

$$u_k = \text{arrival time no } k - \text{arrival time no } k - 1$$

is a sample of the interarrival time sequence (U_k) .

The typical shape for these sequences based on trace data obtained in this study *without silence suppression* is shown in Figure 5, which shows (v_k) and (u_k) for a cut-out section of 200 packets ($1700 \leq k < 1900$), corresponding to four seconds of sending time during one transmission of the pre-recorded voice. To further illustrate such trace data, Figure 6 shows a histogram of the delays (v_k) and Figure 3 a histogram for the interarrival times (u_k) . It can be noted that large values of interarrival times are sometimes followed by very small ones, manifesting that a severely delayed packet forces subsequent packets to arrive end-to-end. The fraction of packets arriving end-to-end corresponds to the height of the leftmost peak in the histogram of Figure 3.

Returning to traces with silence suppression, Figure 7 gives the statistics of the recorded voice signal used. The upper part shows a histogram of the talkspurts and the lower part the corresponding histogram for the non-zero part of the distribution G of the silence intervals X discussed in Section 3. The probability $\alpha = P(X = 0)$ and the expected value $\mu = E(X)$ were estimated to

$$\alpha^* = 0.0456 \quad \mu^* = 25.7171.$$

¹ Available from http://www.sics.se/~ianm/Course/Stochastic_traffic_modelling/Plots/

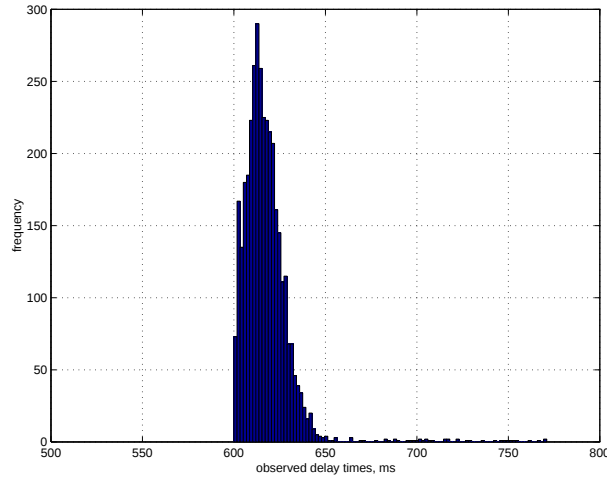


Figure 6: Histogram of the observed delays

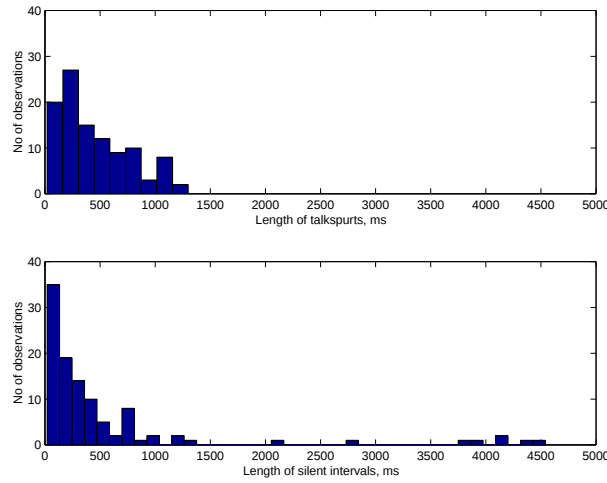


Figure 7: Nature of recorded voice signal

5.2 Numerical estimates

In this section we indicate a few simple numerical techniques that give parameter estimates based on trace data. In principle such methods based on the model presented here can be used for systematic studies of the delays and losses and for comparison of traces sampled in different environments.

Considering first the case of no silence suppression, it was pointed out in Section 4 that given an observed realization $(u_k)_{k=1}^n$ of (U_k) , a point estimate of the packet loss probability p is obtained from (17) (with $\mu = 0$), using

$$p^* = 1 - \frac{20}{\bar{u}}, \quad \bar{u} = \frac{1}{n} \sum_{k=1}^n u_k \text{ ms.}$$

Our measurements gave consistently $\bar{u} \approx 20.002 - 20.005$ ms, indicating loss probabilities of the order 10^{-4} .

Next we look at an experiment where the pre-recorded voice is transmitted at seven different times using silence suppression, and the interarrival times measured at the receiver during each transmission. Table 1 shows the expected silence interval $E(X)$ and the estimated μ from the trace files.

The obtained estimates indicate a systematic bias of the order 0.5 milliseconds in the mean values of the silence suppression intervals. Packet losses do not seem to explain fully the observed deviation. A more comprehensive statistical analysis might reveal the source of this slight mismatch. For the present preliminary investigation we find the numerical estimates satisfactory.

We consider now the problem of estimating the distribution F of packet delays Y given a fixed length sample observation (v_k) of the Markov chain (V_k) for observed delays. One method for this can be based on the steady state analysis

Table 1: Silence Interval Parameters

Trace	$E(X)$	$E(X)-\mu^*$
trace 1	25.7492	0.0321
trace 2	26.2204	0.4639
trace 3	26.2284	0.5113
trace 4	26.2164	0.4993
trace 5	26.2186	0.5015
trace 6	26.2124	0.4953
trace 7	26.2209	0.5038

in Section 2.2. Indeed, rewriting (8) as the simple relation

$$P(V_\infty \leq x) = F(x) \prod_{i=1}^{\infty} F(x+i) = F(x) P(V_\infty \leq x+1)$$

shows that if we let $\bar{F}_V$ denote an empirical distribution function of V , then we obtain an estimate $\bar{F}$ of F by taking

$$\bar{F}(x) = \frac{\bar{F}_V(x)}{\bar{F}_V(x+1)} \quad x \geq 0, \quad (18)$$

where we recall that the variable x is measured in units of 20 ms intervals. An application of this numerical algorithm to the trace data of the previous Figures 5-6 yields an estimated density function for Y as in Figure 8. The numerical scheme

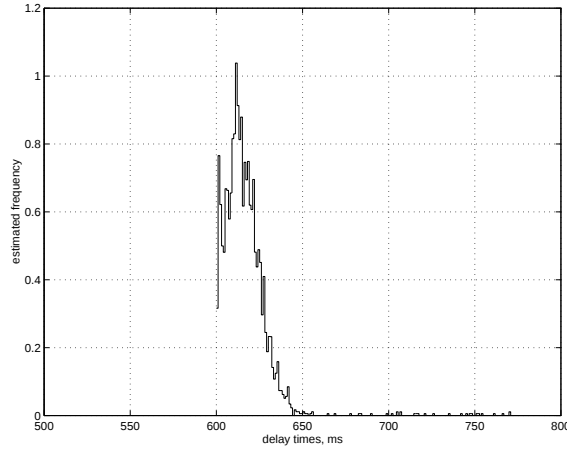


Figure 8: Estimated density of Y

is sensitive for small changes in the data, and it is difficult to draw conclusions on the finer details of the distribution of F . As expected the graph is very similar to that of the observed delays, Figure 6, but with certain differences due to the Markovian dependence structure in the sequence (V_k) as opposed to the independence in (Y_k) . The main difference is the shift towards smaller values for Y in comparison to those of V . This corresponds to the inequality $\bar{F}(x) \geq \bar{F}_V(x)$ valid for all x , which is obvious from (18).

6 Related Work

Many researchers have looked at the needs in terms of buffer size for packet streams characterised by Markov (semi or modulated) behaviour especially in the case for multiplexed sources. Their goal was to derive the waiting time of packets spent in the buffer shown as probability density function of the waiting times. Relatively few however, have looked at the arrival process from a stage of buffers and identifying embedded Markov chains from a single source and concentrating solely on this artifact including both with and without silence suppression. Additionally as far as we know, no-one has used real trace data to enhance and verify their models to the level we show.

Some early analytical work on the buffer size requirements for packetised voice is nicely summarised by Gopal *et al.* [1]. One often cited piece of work is Barberis [2]. As part of this work he assumes the delays experienced by packets

of the same talkspurt are i.i.d according to an exponential distribution $p(t) = \lambda e^{-\lambda t}$ where $1/\lambda$ is the average network delay and standard deviation. M.K. Mehmet Ali *et al.* in their work of buffer requirements [3] model the arrival process as a Bernoulli trial with probability $[1 - F(j, n - j + 1)]$ of the event “no arrival yet” at each interval up to its arrival. The outcome of the trial is represented by the random variable $k(j, n)$:

$$k(j, n) = \begin{cases} 1 & \text{if packet } j \text{ has arrived at or before time } n \\ 0 & \text{otherwise.} \end{cases}$$

Ferrandiz and Lazar in [4] look at the analysis of a real time packet session over a single channel node and compute its performance parameters as a function of their model primitives. They do not use any Markovian assumptions, rather an approach which uses a series of overload and under-load periods. During overload packets are discarded. They derive an admission control scheme based on an average of the packet arrival rate. Van Der Wal *et al.* derive a model for the end to end delay for voice packets in large scale IP networks [5]. Their model includes different factors contributing to the delay but not the arrival process of audio packets per se. The mathematical model described here is also discussed in the forthcoming book [6].

7 Conclusions

We have addressed the problem of modelling the arrival process of a single packet audio stream. Our goal was to present the work by stating the assumptions, building from simple principles and adding complexity to the model in stages. The model can be used to produce packet audio streams with characteristics, at least, quite similar to the particular traces we have obtained. The model is suitable for generating streams both with and without silence suppression, also taking into consideration lost packets. The results of the work can be seen in the density functions of Figure 3 and the model in Figure 8.

The work can be generally applied to research where modelling arriving packet audio streams needs to be performed. Future work includes parameterising the model for specific applications such as where the number of hops needs to be included, or a particular loss rate, or standard coding rate for example. Currently we use typical values but certain modelling tasks may require tuning these parameters. A natural next step is to use the arrival model presented here for evaluation of jitter buffer performance, such as investigating waiting times and possible packet loss in the jitter buffer.

References

- [1] Prabandham M. Gopal, J. W. Wong, and J. C. Majithia, “Analysis of playout strategies for voice transmission using packet switching techniques,” *Performance Evaluation*, vol. 4, no. 1, pp. 11–18, Feb. 1984.
- [2] Giulio Barberis, “Buffer sizing of a packet-voice receiver,” *IEEE Transactions on Communications*, vol. COM-29, no. 2, pp. 152–156, Feb. 1981.
- [3] Mehmet M. K. Ali and C. M. Woodside, “Analysis of re-assembly buffer requirements in a packet voice network,” in *Proceedings of the Conference on Computer Communications (IEEE Infocom)*, Bal Harbour (Miami), Florida, Apr. 1986, IEEE, pp. 233–238.
- [4] Josep M. Ferrandiz and Aurel A. Lazar, “Modeling and admission control of real time packet traffic,” Technical Report Technical Report Number 119-88-47, Center for Telecommunications Research, Columbia University, New York, New York 10027, 1988.
- [5] Walm van der Wal, Mandjes Michel, and Rob Kooij, “End-to-end delay models for interactive services on a large-scale ip network,” IFIP, June 1999.
- [6] Ingemar Kaj, *SIAM Monographs on Mathematical Modeling and Computation*, 2002. To appear., SIAM, 2002.