

SPECTRUM ESTIMATION METHODOLOGY FOR NEXT GENERATION WIRELESS SYSTEMS¹

Tim Irrnich, Bernhard Walke

Chair of Communication Networks, Aachen University, Germany, {tim|walke}@comnets.rwth-aachen.de

Abstract - A new methodology for calculating the spectrum demand for next generation wireless communication systems is presented¹. The methodology considers a cell or cell cluster carrying the traffic of a mix of packet based services, taking the multiplexing gain of packet based traffic and the service-specific QoS parameters mean throughput, mean delay and delay percentile into account. A suitable balance between accuracy and complexity has to be found. Accordingly, two approaches of different complexity are presented. A very simple approach maps the expected offered packet traffic to the required system capacity by exploiting a characteristic of the relation between system load and system throughput that is common to all packet-based systems. A second approach is based on the analysis of a queuing system. The required system capacity is calculated for services with mean throughput and mean delay requirements. Using an additional iterative procedure also delay percentile requirements can be taken into account. The queuing approach is more complex but also significantly more accurate than the simple approach.

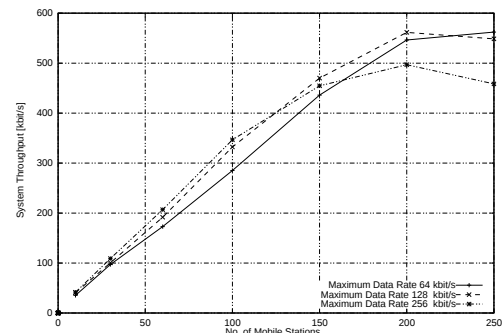
Keywords - Spectrum estimation, 4G, Next Generation Wireless Systems, M/G/1, NONPRE, priorities, mean waiting time, mean IP delay, delay percentile

I. INTRODUCTION

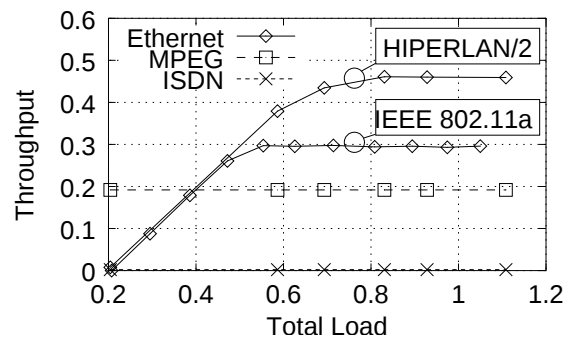
Current methodologies for spectrum estimation calculate the system capacity needed to fulfill the offered traffic's *Quality of Service* (QoS) requirements based on the Erlang theory [1]. Owing to significant advances in technology [2], traffic models and user demands, the Erlang theory is not sufficient to model the behavior of current and future wireless systems, and a new methodology is needed to support future decisions in the international spectrum regulatory process [3]. A key aspect that has to be addressed by a new methodology is the modelling of packet switched (PS) services and systems [4]. For packet traffic a new methodology has to offer a well-balanced trade-off between accuracy and complexity.

This paper presents a new methodology for calculating the spectrum requirements for packet traffic. Two alternative approaches to calculate the required system capacity for packet traffic under QoS constraints are introduced.

¹This work has been funded by the German Federal Ministry for Research and Education (BMBF) in the Multifunk project and by the IST project WINNER



(a) UMTS (source: [2], p. 540)



(b) Wireless LANs (source: [2], p. 850)

Fig. 1

System throughput vs offered traffic shows the same behavior for different Radio Access Technologies

The first approach is based on the observation that the relation between offered traffic/system load and system throughput of a large number of wireless systems basically shows the same behavior. With increasing traffic load up to the saturation point the system throughput is equal to the offered traffic. At the saturation point additional offered traffic does no longer lead to higher system throughput. In Fig. 1 this behavior is visualized for UMTS and for IEEE 802.11 and HIPERLAN/2.

The second approach models a wireless system as a M/G/1 queue with non-preemptive priorities and *First-Come-First-*

Serve (FCFS) scheduling within the packets of equal priority. Non-preemptive priority means that upon arrival of a job with higher priority than the current job, the current is not interrupted, but completed before service of the newly arrived higher priority job is started. For each service type considered one priority level can be used. Based on a certain mean delay required for the packets of each priority level, the required service channel data rate (i.e. the system capacity) is calculated. An additional iterative procedure allows to adjust the capacity to fulfill a certain delay percentile requirement for each service type.

II. TYPES OF COMMUNICATION

With respect to the impact on spectrum demand it is necessary to distinguish the following types of communication:

- A) Human-to-Human (e.g. Speech): A certain percentage, denoted by a , of the traffic that belongs to this type of communication has both ends of the connection located in the same cell cluster. Accordingly, the data has to be transmitted twice over the air interface, and the spectrum requirement for this traffic has to be multiplied with a factor $(1 + \frac{a}{100})$.
- B) Human-to-Machine (e.g. WWW): All traffic is only transmitted once over the air interface.
- C) Machine-to-Machine (e.g. SMS, file sharing, peer-to-peer applications): The peer device can be a fixed device or a mobile device. To account for the mobile devices the spectrum requirement for this type of communication has to be multiplied with a factor $(1 + \frac{b}{100})$.
- D) Broadcast: All traffic has to be transmitted only once in each cell of the cluster. Due to the fact that this type of communication is point-to-multipoint and only occupies resources in the downlink direction, it can be considered as one user per cell using a corresponding traffic class. Since there might be d cells in a cluster the spectrum requirement for this type has to be multiplied with a factor $(1 + \frac{d}{100})$.

To simplify the formulae presented in the following, these factors have not been included throughout the rest of this paper.

III. INPUT PARAMETERS FOR SPECTRUM CALCULATION

The values of the following parameters are needed for both approaches. They are assumed to be known under high load of a future system, e.g. for the busy hour:

- Number of different service types N
- Aggregate offered traffic of all users per service type T_n on IP layer per cell or sector (unit $\frac{\text{bps}}{\text{cell}}$).
- Area Spectral Efficiency (ASE) (unit $\frac{\text{bit/s}}{\text{Hz} \cdot \text{m}^2}$), denoted by η . The ASE is assumed to be obtained regarding a fully loaded system.
- Area A covered by the cell or sector regarded for determination of the offered traffic.

- The channel bandwidth B needs to be specified for determination of the ASE. The resulting spectrum demand is an integer multiple of B .

IV. TRAFFIC MODELLING

We assume that the distribution of the traffic to different alternatively available *Radio Access Technologies* (RAT) and to cell layers inside each RAT has already been taken into account. The spectrum demand *per RAT* can be determined by applying our methodology separately to each cell layer and.

For each service type the offered traffic is considered on IP layer. The traffic is characterized by IP packet inter-arrival time and packet size distribution.

For specific applications or codecs packet inter-arrival time and packet size distributions can be obtained from measurements or simulations. For each service type the analysis of the M/G/1 queue requires the mean arrival rate of IP packets and the first three moments of the packet size distribution as input parameters.

V. REQUIRED SYSTEM CAPACITY

In this step the required system capacity needed serve the offered base traffic while fulfilling the QoS requirements is determined.

For RATs using TDD with flexible TX/RX border the offered base traffic contains the traffic of both, uplink and downlink direction. Otherwise the required system capacity has to be calculated separately for uplink and downlink. In order to improve readability the calculation is shown exemplarily for one direction, or the TDD case with flexible RX/TX border, respectively.

A. Approach 1: Empirical QoS Factor

A schematic view of the typical relation between offered traffic and aggregate system throughput of a wireless system is shown in Fig. 2. The maximum of the system throughput denotes the capacity that is theoretically available. Since the offered traffic on each layer consists of original user traffic and overhead added from higher layers, the capacity that is needed on physical layer consists of a part that is needed for actual user traffic and a part that is needed for system overhead.

It is widely known that the packets in a saturated system encounter infinitely long delay. In order to fulfill the delay requirements of delay-sensitive service types only some percentage of the system capacity can be utilized. The fraction of the theoretical system capacity that can be utilized to be able to meet the QoS requirements of the regarded traffic classes is called the *Effective Capacity*.

The gap between the theoretical system capacity C_t and the effective capacity C_e is called the *QoS Reserve*, the corresponding relation C_e/C_t is called the *QoS Factor* κ .

$$C_e = \kappa \cdot C_t \Leftrightarrow C_t = \frac{C_e}{\kappa}; \quad \kappa \leq 1. \quad (1)$$

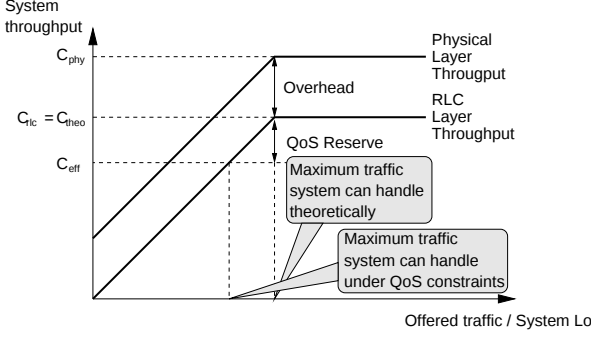


Fig. 2

The effective capacity with respect to QoS constraints is lower than the theoretical capacity

For a given set of QoS parameters defined for a single service type κ can be empirically determined from analytical and/or simulation models. A gross estimate could be $\kappa = 0.75$.

In addition to the input parameters listed in Sec. III, this approach requires the following parameters:

- QoS factor κ .
- Protocol overhead ratio O_n .
- Packet Error Ratio (PER) R .

Using Eq. (1) we determine the theoretical capacity that is needed to fulfill the QoS requirements of service type n . In order to calculate the spectrum bandwidth required for a system to be able to provide the required theoretical capacity, the required capacity on PHY level, C_p has to be determined. Packet errors and protocol overhead are taken into account, see Fig. 2. Assuming an ideal SR-ARQ protocol the required physical layer capacity for service type n can be calculated to

$$C_{p,n} = \frac{C_{t,n}}{(1-R)(1-O_n)} \quad (2)$$

$$= \frac{C_{e,n}/\kappa_n}{(1-R)(1-O_n)} = \frac{T_n/\kappa_n}{(1-R)(1-O_n)}.$$

The existence of one unique PER value for the regarded area A is assumed. The PER mainly depends on the distance between Base Station and Mobile Station and therefore is not constant in the regarded area. In order to obtain a conservative approximation, we select the PER value that corresponds to the longest distance between BS and MS in the regarded area A .

O_n denotes the protocol overhead percentage, i.e., the ratio between user data at the RLC layer and data to be transmitted on PHY layer. For some RATs the protocol overhead might depend on the service type due to some kind of reconfigurability of the protocol stack.

The total physical layer capacity required to serve all

service types according to their QoS constraints is

$$C_p = \sum_{n=1}^{N_s} \frac{T_n/\kappa_n}{(1-R)(1-O_n)}. \quad (3)$$

B. Approach 2: Queuing model

We select the M/G/1/FCFS queue with non-preemptive priorities (also known as head-of-the-line queuing system) as a model for a cell that has to carry the traffic of a number of different service types, each having different requirements of the mean packet delay. For this queue a well-established analysis is available [5], [6].

One job served by the queuing system is defined as one IP packet. By using non-preemptive priorities it is assumed that each IP packet is completely served before the current radio resource allocation is changed. This is a valid assumption, because in many cases interrupting service for an IP packet in service causes loss of the capacity spent already for that packet. In data communication systems (like the Internet) an ongoing transmission of a limited size data unit (a packet) is never interrupted in favor of another one.

A RAT is modeled here as having a single packet channel only, independent of the number of channels used in parallel in a real RAT, since there is no trunking gain possible when multiplexing packets buffered in a queue to be transmitted via one or more parallel channels. Some minor overhead resulting from fragmentation and padding when using multiple parallel medium bit rate channels instead of one equal capacity high bit rate channel is neglected.

1) *Services with mean delay requirements:* In addition to the parameters listed in Sec. III, for this approach the following input parameters are needed:

- Mean and second moment of the IP packet size distribution of service type n , denoted by s_n and $s_n^{(2)}$, where $s_n^{(2)} = s_n^2 + \sigma_n^2$ and σ_n^2 denotes the variance.
- The required mean delay D_n of each service type
- The priority ranking of all service types n with $n = 1, 2, \dots, N_{ps}$. In the following formulas it is assumed that the service type $n = 1$ has the highest priority, i.e. IP packets of service type $n = 1$ are served first. The service type $n = N_{ps}$ has the lowest priority.
- A correction factor $\zeta_n > 1$ in order to take the underestimation of the calculated required system capacity into account, that is caused by the fact that packet arrivals are not completely independent, in contrast to the assumption of the queuing model used.

According to Cobham in a M/G/1/FCFS-NONPRE system the mean waiting time of priority n is

$$W_n = \frac{\lambda_{\leq N} \beta_{\leq N}^{(2)}}{2(1-\rho_{\leq n})(1-\rho_{\leq (n-1)})}, \quad (4)$$

where

$$\lambda_{\leq N} = \sum_{i=1}^N \lambda_i \quad (5)$$

is the aggregated arrival rate of all priority levels,

$$\beta_{\leq n}^{(r)} = \sum_{i=1}^n \frac{\lambda_i}{\lambda_{\leq n}} \beta_i^{(r)} \quad (6)$$

is the r th moment of the weighted service time distribution of all priority levels and

$$\rho_{\leq n} = \sum_{i=1}^n \rho_i = \sum_{i=1}^n \lambda_i \beta_i \quad (7)$$

is the aggregate system load of the priorities less than and equal to level n .

The resulting IP packet arrival rate λ_n (unit: packets/s) of service type n is obtained by dividing the offered base traffic by the mean packet size,

$$\lambda_n = \frac{T_n}{s_n}. \quad (8)$$

We now express the mean IP packet delay for packets of priority n as the sum of mean waiting time W_n and service time β_n ,

$$D_n = W_n + \beta_n \quad (9)$$

and express the service time for an IP packet T_s as packet size S divided by the service channel's data rate C ,

$$T_s = \frac{S}{C}. \quad (10)$$

Accordingly the mean service duration of an IP packet of service type n is equal to the mean IP packet size of service type n divided by the service channel's data rate,

$$\beta_n = \frac{s_n}{C}. \quad (11)$$

Since this is equivalent to a linear transformation of one random variable to another random variable, we obtain the following relation between the second moments of the service time distribution and the packet size distribution:

$$\beta_n^{(2)} = \sum_{\forall t} p(T_s = t) \cdot t^2 = \sum_{\forall x} p(S = x) \cdot \left(\frac{x}{C}\right)^2 = \frac{s_n^{(2)}}{C^2}. \quad (12)$$

Accordingly the relation between $\beta_{\leq N}^{(2)}$ and $s_{\leq N}^{(2)}$ is

$$\beta_{\leq N}^{(2)} = \frac{s_{\leq N}^{(2)}}{C_n^2}, \quad s_{\leq N}^{(2)} = \sum_{i=1}^N \frac{\lambda_i}{\lambda_{\leq N}} \cdot s_n^{(2)} \quad (13)$$

Inserting Eqs. (11), (12) and (13) into Eq. (5), the system capacity C_n that is needed to obtain the mean delay required by service type n can be calculated from

$$D_n - \frac{s_n}{C_n} = \frac{\lambda_{\leq N} s_{\leq N}^{(2)}}{2 \left(C_n \sum_{i=1}^n \lambda_i s_i \right) \left(C_n \sum_{i=1}^{n-1} \lambda_i s_i \right)}. \quad (14)$$

Since the required system capacity calculated this way only fulfills the mean delay requirements of the service

type n , Eq. (5) has to be solved for each priority level in order to determine the most demanding priority level. The computational efforts can be reduced by pooling multiple service types with identical QoS requirements into one priority level.

The priority level which requires the highest capacity then denotes the required system capacity for all service types, since for the case that the QoS requirements of the most demanding service type are fulfilled, the requirements of the other service types are over-fulfilled. The mean throughput requirement is automatically fulfilled for all service types when the aggregate system load (see Eq. (7)) is less than one, i.e. the queue is operated in a stable state.

After the C_n that is needed to achieve the required mean delay for service type n is determined, we apply a correction factor that accounts for some additional capacity that is required owing to correlated packet arrivals, which have been neglected up to this point by assuming the independent property of the arrival process,

$$C'_n = \zeta_n \cdot C_n. \quad (15)$$

The correction factor can be empirically derived by validating the results of the M/G/1/FCFS-NONPRE queue with results for more complicated queuing models, simulation or measurement results.

The overall required system capacity is determined by the most demanding service type,

$$C = \max(C_n, \forall n). \quad (16)$$

2) *Services with delay percentile requirements:* In addition to the mean delay to be met, some service types might require a certain percentage of all packets to have a delay below some certain value. This kind of requirement is called a percentile requirement.

In this section we present an iterative algorithm that allows to adjust the system capacity for service types that in addition to the mean delay require some percentile of the packet delay CDF.

Additional input parameter needed for this step are:

- The third moment of the packet size distribution $s_n^{(3)}$.
- The typical packet size of a short and a long IP packet for service type n , denoted by $x_{short,n}$ and $x_{long,n}$, respectively.
- The required percentile (e.g. the 95% percentile) for each service type.
- The system capacity C_n that was determined according to Sec. V-B to meet the mean delay requirement of service type n .

We consider the waiting time and service duration CDFs to check if a given delay percentile of a service type under consideration is met. If not, the capacity that was determined to meet the requirement for the mean delay will have to be adjusted (i.e. increased).

First we determine the second moment of the waiting time distribution. According to [6] in a M/G/1/FCFS-NONPRE

queue the second moment of the waiting time distribution can be obtained from

$$W_n^{(2)} = \frac{\lambda_{\leq N} s_{\leq N}^{(3)}}{3 \left(C_n \sum_{i=1}^n \lambda_i s_i \right) \left(C_n \sum_{i=1}^{n-1} \lambda_i s_i \right)^2} + \frac{\lambda_{\leq N} s_{\leq N}^{(2)} \lambda_{\leq n} s_{\leq n}^{(2)}}{2 \left(C_n \sum_{i=1}^n \lambda_i s_i \right)^2 \left(C_n \sum_{i=1}^{n-1} \lambda_i s_i \right)^2} + \frac{\lambda_{\leq N} s_{\leq N}^{(2)} \lambda_{\leq n} s_{\leq n}^{(2)}}{2 \left(C_n \sum_{i=1}^n \lambda_i s_i \right) \left(C_n \sum_{i=1}^{n-1} \lambda_i s_i \right)^3}, \quad (17)$$

where

$$s_{\leq N}^{(3)} = \sum_{i=1}^N \frac{\lambda_i}{\lambda_{\leq N}} s_i^3 \quad (18)$$

denotes third central moment of the weighted common packet size distribution of all priority levels together. In Eq. (17) the first three moments of the weighted common service time distribution have been substituted in analogy to Eq. (13).

Now it is possible to approximate the waiting time CDF by an analytical function that matches the first two moments of the waiting time CDF. From queuing theory is known that for M/G/1 types of systems the tail (i.e. for large delays) of the delay CDF is exponential. Thus, a suitable approximation of the waiting time distribution is the degenerated hyper-exponential distribution of second order. Denoting the waiting time random variable by T_w this CDF is

$$P_n(T_w \leq t) = 1 - \rho_n e^{-\gamma_n t} \quad (19)$$

with

$$\gamma_n = 2 \frac{W_n}{W_N^{(2)}} \quad \text{and} \quad \rho_n = 2 \frac{W_n^2}{W_N^{(2)}}. \quad (20)$$

The corresponding PDF can be obtained by derivation of (19) to t ,

$$p_n(T_w = t) = \frac{d}{dt} (1 - \rho_n e^{-\gamma_n t}) = \rho_n \gamma_n e^{-\gamma_n t}. \quad (21)$$

Owing to the internal correlation structure of the arrival process the waiting time CDF could alternatively be approximated using the Pareto distribution. In this case the convolution integral for determination of the IP packet delay *Probability Density Function* (PDF) might no longer have a closed form solution, and a numerical algorithm would be required to solve Eq. (28). In this case it would be better to execute the convolution operation by multiplying the Laplace-Stieltjes transforms (LST) of waiting time and service time approximation PDFs. If the closed inverse transformation of the result is not possible in this case, the delay PDF could again be approximated by a number of moments, since the r th moment of a distribution can be

determined by from the r th derivative of the LST, evaluated at $s = 0$.

To obtain the delay PDF is further required to approximate the service time PDF by an analytical function that matches the first two moments of the service time distribution. It is widely known that an IP packet size distribution is dominated by two types of packets, short (i.e. nearly empty except for header information) and long packets (according to the IP *Maximum Transfer Unit* (MTU)). Thus, a suitable approximation of the service time PDF is a function that consists of two Dirac impulses, located at the service time for a typical short and long packet of service type n , respectively. These service times are denoted by $t_{short,n}$ and $t_{long,n}$. The corresponding PDF is

$$p_n(T_s = t) = p_{1,n} \cdot \delta(t - t_{short,n}) + p_{2,n} \cdot \delta(t - t_{long,n}) \quad (22)$$

$$\text{with } p_{1,n} + p_{2,n} = 1,$$

where $\delta(t)$ denotes the Dirac impulse. $p_{1,n}$ and $p_{2,n}$ can be determined from

$$\begin{aligned} \beta_n &= p_{1,n} \cdot t_{short,n} + p_{2,n} \cdot t_{long,n} \\ &= p_{1,n} \cdot \frac{x_{short,n}}{C_n} + p_{2,n} \cdot \frac{x_{long,n}}{C_n} \end{aligned} \quad (23)$$

and

$$\begin{aligned} \beta_n^{(2)} &= p_{1,n} \cdot t_{short,n}^2 + p_{2,n} \cdot t_{long,n}^2 \\ &= p_{1,n} \cdot \frac{x_{short,n}^2}{C_n^2} + p_{2,n} \cdot \frac{x_{long,n}^2}{C_n^2}, \end{aligned} \quad (24)$$

leading to the following expressions

$$p_{1,n} = \frac{\beta_n \left(\frac{x_{short,n}}{x_{long,n}} - x_{short,n} - x_{long,n} \right) + \beta_n^{(2)} C_n}{\frac{x_{short,n}}{C_n} \left(\frac{x_{short,n}}{x_{long,n}} - x_{long,n} \right)} \quad (25)$$

and

$$p_{2,n} = C_n \frac{\beta_n x_{short,n} - \beta_n^{(2)} C_n}{x_{short,n} - x_{long,n}^2}. \quad (26)$$

After approximating both, waiting time distribution and service time distribution with analytical functions as described above and denoting the IP packet delay random variable by T_d we obtain an approximation for the IP packet delay PDF by convolution of the service time and waiting time PDFs,

$$p_n(T_d = \tau) = \int_{-\infty}^{+\infty} p_n(T_w = t) \cdot p_n(T_s = (\tau - t)) dt, \quad (27)$$

and the IP packet delay CDF by integration of the PDF,

$$P_n(T_d \leq \tau) = \int_{-\infty}^{+\tau} p_n(T_d = t) dt. \quad (28)$$

Inverting the resulting IP packet delay CDF allows an immediate check whether the percentile specified for a given service type can be met by the capacity C_n calculated on basis of the mean delay requirements.

If the specified percentile is not be met under the capacity calculated before, the capacity will have to be increased, say by 10%,

$$C_n' = C_n \cdot 1.1, \quad (29)$$

and the delay CDF calculated anew to make another trial to meet the requirement. This procedure is repeated until the delay percentile requirement is met.

As described above the overall required system capacity is determined by the most demanding service type,

$$C = \max(C_n', \forall n). \quad (30)$$

VI. REQUIRED SPECTRUM PER CELL OR SECTOR

The required spectrum is obtained by application of the ASE and rounding up to the next higher integer number of the channel bandwidth that was assumed for determination of the ASE value,

$$F = B \cdot \left\lceil \frac{1}{B} \cdot \frac{C}{\eta \cdot A} \right\rceil \quad (31)$$

or

$$F = B \cdot \left\lceil \frac{1}{B} \cdot \frac{C_{up} + C_{down}}{\eta \cdot A} \right\rceil, \quad (32)$$

respectively.

Since the dimensioning procedure described above results in a system carrying the maximum possible amount of traffic so that it is possible to fulfill the QoS requirements of all service types, the spectral efficiency has to be valid for a fully loaded system.

There might be different systems A and B that reach slightly different spectral efficiencies F_A and F_B under certain comparable conditions, but have substantially different cost of infrastructure, user equipment and operations. In such a case it appears natural to provide low value service types in the low cost system, since these otherwise would not be accepted from the customer. A consequence of this is that the traffic amount would be much larger than observable when the higher cost system would offer the same services. This could result in higher spectrum requirements for the low cost system compared to the high cost system.

VII. CONCLUSION

This paper presents a novel approach to calculate the spectrum demand for a future cellular mobile communication system.

Two alternative approaches, each representing a different trade-off between accuracy and complexity are presented. Based on well-established results from traffic theory, these

new concepts are robust and technology-neutral. The multiplexing gain resulting from radio resource sharing between packet-based services and service-specific QoS constraints are taken into account. QoS is considered in terms of mean throughput, mean delay and delay percentiles.

Considering the accuracy that can be expected for input parameters and underlying assumptions like the estimation of the amount of traffic and the market penetration of future applications and the spectral efficiency of future systems, the accuracy of the presented calculations is sufficient.

There are still a number of open issues to be addressed in future research:

- Multihop communication resulting from relay-based deployment concepts is widely agreed to be an essential part of 4G systems. It is expected that multihop will significantly increase spectral efficiency.
- The amount of spectrum needed also depends on the grade of mobility (e.g. in terms of resources reserved for call handover). Up to now, the influence of user mobility is not considered. A straightforward way to do so would be via the spectral efficiency.

ACKNOWLEDGMENT

The authors would like to thank their partners in the IST-WINNER project for fruitful discussions during the preparation of this paper. Especially we would like to thank Joerg Huschke at Ericsson GmbH Aachen and Antti Lappetelainen at Nokia Research Center Helsinki for their valuable comments.

REFERENCES

- [1] ITU-R, "Methodology for Calculation of the IMT-2000 Terrestrial Spectrum Requirements," ITU, Recommendation M.1390, 1999.
- [2] B. Walke, *Mobile Radio Networks - Networking, Protocols and Traffic Performance*, 2nd ed. John Wiley & Sons, 2002.
- [3] ITU-R, "Framework and Overall Objectives for the Future Development of IMT-2000 and Systems Beyond IMT-2000," ITU, Recommendation M.1645, 2003.
- [4] F. R. of Germany, "One example of a methodology to calculate spectrum requirements," ITU-R, Working Party 8F Doc. 8F/881-E, 2003.
- [5] A. Cobham, "Priority assignments in waiting line problems," *Operations Research*, vol. 2, pp. 70–76, 1954.
- [6] M. Schwartz, *Computer Communication Network Design and Analysis*. Prentice-Hall, 1977.