

Foundations and Trends® in Networking
Vol. 12, No. 1-2 (2017) 1–161
© 2017 T. de Cola, A. Ginesi, G. Giambene,
G. C. Polyzos, V. C. Siris, N. Fotiou and
Y. Thomas
DOI: 10.1561/13000000046



Network and Protocol Architectures for Future Satellite Systems

Tomaso de Cola
German Aerospace Center
tomaso.decola@dlr.de

Alberto Ginesi
European Space Agency
alberto.ginesi@esa.int

Giovanni Giambene
University of Siena
giambene@unisi.it

George C. Polyzos
Athens University of Economics and Business
polyzos@aueb.gr

Vasilios A. Siris
Athens University of Economics and Business
vsiris@aueb.gr

Nikos Fotiou
Athens University of Economics and Business
fotiou@aueb.gr

Yiannis Thomas
Athens University of Economics and Business
thomasi@aueb.gr

Contents

1	Introduction	8
1.1	Modern Satellite Systems	8
1.2	Overall Framework	10
2	Future Trends in Satellite Communications	12
2.1	High Throughput Satellite (HTS) systems	12
2.2	Non-GEO Satellite Systems	45
2.3	Networking Challenges	61
3	Networking Solutions for High Throughput Satellite Systems	73
3.1	MPTCP	73
3.2	Network Coding and its Applications to Satellite Networking	93
4	Future Trends in Satellite Networking	119
4.1	SatCom Use Cases and Applications	120
4.2	Future SatCom based on ICN architectures	132
	Bibliography	152

Abstract

Since their conception, satellite communications have been regarded as a promising tool for all environments where the terrestrial infrastructure is limited in capacity or to take advantage of the multicasting/broadcasting capabilities inherent in satellite technology. Recent advances have seen satellite technology mature to a more prominent role in the telecommunications domain. In particular, the design of novel satellite payload concepts for Geostationary (GEO) satellite platforms, as well as renewed interest in Low Earth Orbit (LEO) satellite constellations have made the integration of satellite and terrestrial networks almost compulsory to ensure new services meet the requirements for high user-rate and quality of experience that could not be achieved using either of the two technologies independently. From this viewpoint, convergence of satellite and terrestrial technologies also requires considering the most recent trends in networking, with special attention being paid to the potential new architectures that have been recently proposed in the framework of Future Internet.

This monograph explores the main components of the scenarios above, putting particular emphasis on the networking aspects. To this end, novel protocols such as Multi Path TCP (MPTCP) and networking trends such as Information Centric Networking (ICN) are explored by demonstrating their applicability in some scenarios that deploy both satellite and terrestrial segments. Particular attention is given to smart gateway diversity schemes which advocate the use of sophisticated multi-path transmission schemes to exploit the multi-homing features offered by present day devices. The second part of the monograph is dedicated to content-based networking, which is becoming increasingly popular driven by the pervasiveness of the Internet in everyday life. In this regard, applications to satellite communications are illustrated and the technical challenges to be further addressed are highlighted.

List of Abbreviations and Acronyms

ACK	Acknowledgment
ACM	Adaptive Coding and Modulation
AeNB	Aerial eNB
AIDA	Agile Integrated Downconverter Assembly
AIMD	Additive Increase Multiplicative Decrease
AMR	Automatic Meter Reading
AP	Access Provider
ARQ	Automatic Repeat reQuest
AS	Autonomous System
BDP	Bandwidth-Delay Product
BER	Bit Error Rate
BFN	Beam Forming Network
BH	Beam Hopping
BIC	Binary Increase Congestion control

BSM	Broadband Satellite Multimedia
CCSDS	Consultative Committee for Space Data Systems
CDN	Content Delivery Networks
CP	Content Provider
CR	Content Router
CRA	Contention Resolution ALOHA
CRC	Cyclic Redundancy Check
CRDSA	Contention Resolution Diversity Slotted ALOHA
cwnd	congestion window
DAMA	Demand Assignment Multiple Access
DPI	Deep Packet Inspection
DRA	Direct Radiating Array
DSA	Diversity Slotted ALOHA
DTN	Delay/Disruption Tolerant Network
EIRP	Effective Isotropic Radiated Power
EPC	Evolved Packet Core
ESA	European Space Agency
FAFR	Focal Array Fed Reflector
FCFS	First Come, First Served
FEC	Forward Error Correction
FIFO	First In, First Out
FN	Forwarding Node

FTP	File Transfer Protocol
GEO	Geostationary Orbit
GFP	Generic Flexible Payload
GW	Gateway
HAP	High Altitude Platform
HAP	High-Altitude Platform
HL-BFN	High level BFN
HPA	High Power Amplifier
HTS	High Throughput Satellite
HTS	High Throughput Systems
HTTP	Hypertext Transfer Protocol
ICN	Information Centric Networking
IMUX	Input Multiplexer
IoT	Internet of Things
IP	Internet Protocol
IRIS	IP Routing in Space
IRSA	Irregular Repetition Slotted ALOHA
ISL	Inter-Satellite Link
ISL	Inter-Satellite Links
ISP	Internet Service Provider
LEO	Low Earth Orbit
LFU	Least Frequently Used

LL-BFN	Low level BFN
LNA	Low Noise Amplifier
LRU	Least Recently Used
M2M	Machine-to-Machine
MAC	Media Access Control
MFPB	Multi Feed Per Beam
MPA	Multi-Port Amplifier
MPLS	Multi-Protocol Label Switching
MSS	Maximum Segment Size
NACK	Negative Acknowledgment
NASA	National Aeronautics and Space Administration
NC	No Caching
NCC	Network Control Center
NDN	Named Data Networking
NFV	Network Function Virtualization
NMC	Network Management Center
NRS	Name Resolution Service
OBP	On-Board Processor
PBR	Policy-Based Routing
PEP	Performance Enhancing Proxy
PER	Packet Erasure Rate
PER	Packet Error Rate

PLA	Packet Level Authentication
PLMU	Portable Land Mobile Unit
PLR	Packet Loss Rate
PSI	Publish-Subscribe Internetworking
QoE	Quality of Experience
QoS	Quality of Service
RA	Random Access
RASE	Routing and Switching Equipment
RENE	Rendezvous Network
RLNC	Random Linear Network Coding
RN	Rendezvous Nodes
RTT	Round-Trip Time
SA	Slotted ALOHA
SACK	Selective Acknowledgment
SCACE	Single Channel Agile Converter Equipment
SC-ARQ	Selective-Coded ARQ
SCPS-TP	Space Communications Protocol Specifications - Transport Protocol
SDN	Software Defined Networking
SFPB	Single Feed Per Beam
SIC	Successive Interference Cancellation
SNACK	Selective Negative Acknowledgment
SNO	Satellite Network Operator

SR-ARQ	Selective-Repeat ARQ
SSPA	Solid State Power Amplifier
STP	Satellite Transport Protocol
SVNO	Satellite Virtual network Operator
TCP	Transmission Control Protocol
TCP	Transmission Control Protocol
TM	Topology Manager
TP	Transit Provider
TWTA	Travelling Wave Tube
UAV	Unmanned Aerial Vehicle
V2I	Vehicle-to-Infrastructure
V2V	Vehicle-to-Vehicle
VANET	Vehicular Ad-Hoc Network

1

Introduction

1.1 Modern Satellite Systems

Ever since the inspirational and visionary article from Arthur L. Clark in 1945¹, satellite communications have become more and more part of our everyday life, as they counted a large number of applications such as TV broadcasting, Earth observation, navigation-assisted vehicle, support to disaster situations, just to cite a few. As a result of the increasing number of applications, the satellite academic and industrial community has put quite some effort in developing new platforms able to offer more capacity, so as to enable richer services. From this standpoint, it is also worthwhile to mention the proliferation of communication standards developed to ensure interoperability between different satellite systems, such as those elaborated in DVB and then ETSI standardisation fora.

In the continuous technological progress observed in the last 20 years, a prominent role has been played by the communication paradigm switch from single-beam to multi-beam, in order to provide

¹“Extra-Terrestrial Relays – Can Rocket Stations Give Worldwide Radio Coverage?”, *Wireless World*, October 1945.

larger data-rates, though at cost of increased interference to be contrasted by suitable mitigation techniques. This revolution has led to rethinking of the overall satellite system, for what concerns both the terrestrial and the space segment. As to the latter, classical bent-pipe satellites have been more often accompanied by on-board processing ones, thus broadening the optimisation space to be considered during the system design. In particular, the advent of satellite payload flexible in power, frequency or time (beam-hopping) introduced a new dimension in the resource allocation problem across the entire satellite system, hence facilitating to more efficiently meet the capacity requests of users.

Another key technological advance has been offered by the introduction of LEO constellations in early 2000's that initially turned out be unsuccessful and eventually becoming again an appealing concept as proven by the recent launch of *mega-constellations*, supposed to be able to better serve users with larger data rates and lower access latency, thus possibly resulting a direct competitor to terrestrial technologies. In this perspective, the advent of free space laser optics too signed an important step in revolutionizing the design of future satellite systems, in that they can offer much larger data rates than those available with radio-frequency counterpart, although the performance of the former can be severely hampered by adverse conditions such as clouds.

In spite of the ever-increasing effort made by the satellite community to evolve the operational concept of satellite communication, it is however immediate to grasp that satellite technology cannot be ultimate vehicle to support telecommunications in all its forms. On the contrary, Internet has been typically transported over terrestrial infrastructures and its predominance will even increase, taking also advantage of the increasing penetration of mobile devices in everyday life. Nevertheless, the ideal compromise between the two competing worlds consists in the convergence in a unique ecosystem therefore able to meet all users' demands on a full anytime-anywhere scale. To make the integration exercise meaningful for both worlds, satellite systems has undergone important enhancements from a communication viewpoint, aimed at increasing the overall offered capacity, as testified by the ex-

perimentation of Extra High Frequency (EHF) frequency bands and the related use of diversity techniques to efficiently support gateway handover events and still to attain very high level of system availability. Further to this, new networking paradigms have been explored to let the satellite technology become an appealing candidate for integration with terrestrial network. In the perspective, an important role is also being played by the current reshaping of Internet delivery infrastructures that are more and more tailored around the content rather than the traditional *source-destination* philosophy. From this standpoint, the promotion of Information Centric Networking paradigms represents an important shift in the networking paradigm used so far and also introduces some important features to ease integration between heterogeneous technologies.

All in all, these are the main components that are considered instrumental to develop a more modern vision of satellite systems, which are destined to seamlessly integrate with terrestrial infrastructure in the near future.

1.2 Overall Framework

This monograph surveys the most recent advances in satellite communication technology, putting special emphasis on the networking concepts that are expected to enable seamless integration of satellite and terrestrial segments. In this view, it guides the readers along a path ideally connecting the current trends in satellite payloads design and the related implications in the design of resource allocation schemes with the modern protocol architectures that have emerged during the last years in the terrestrial domain. The logical decomposition of this picture therefore consists in three main elements to which specific sections are reserved, starting from a system view of satellite environments to conclude with an architectural perspective. In this light, the monograph is conveniently structured as follows:

- Section 2 illustrates the main concepts behind the design of flexible and beam hopping payloads, giving also insights into how more efficient resource allocations should be implemented. The

overall discussion provides a system view analysis of satellite systems, providing a possible outlook on the design of next generation satellite systems and consequent enabling of new services.

- Section 3 approaches the trend of network convergence for satellite and terrestrial segments, delving the potentials of multi-path communication protocols. In this respect, overview of the Multi Path TCP protocol (MPTCP) is given and its application combined to networking coding in heterogeneous terrestrial-satellite links is illustrated.
- Section 4 is the natural follow-up of the discussion about integrated satellite and terrestrial network given in Section 3, here giving an architectural perspective. In particular, the recently conceived concept of Information Centric Networking (ICN) is applied to illustrate the advantages in terms of seamless network integration offered by some of its features.

2

Future Trends in Satellite Communications

This chapter presents the main trends in satellite communications, taking as reference the paradigm of High Throughput Satellite (HTS) systems, putting some emphasis on the design choices available for satellite payloads. Later on, other satellite systems configurations are taken into account, so as to introduce the general topic of network convergence between satellite and terrestrial segments, which is expected to be the future of telecommunication systems. In this respect, the most important trends coming up from a networking perspective are shortly outlined, some of them (e.g., MPTCP and network coding) being further illustrated in the following chapters.

2.1 High Throughput Satellite (HTS) systems

2.1.1 Generalities

The ever-increasing demand for high quality and data rate services is going to revolution the telecommunication world, so as to provide users with unprecedented experience in terms of delivered bit rate[19]. On the other hand, digital divide is still affecting several worldwide regions, because of the lack of effective terrestrial infrastructures or

challenging operations conditions due to territory morphology (e.g., mountains, deserts, etc.). From this perspective, the role that satellite is expected to play will be prominent in complementing the terrestrial infrastructure so as to provide all users with very high data rates (e.g., > 50 Mbps as targeted in Europe for years 2020+). To this end, the concept of high throughput satellite (and likewise very high and more recently ultra high) systems has been conceived [36], under the general design mission of achieving Terabit/s capacity [18]. Achieving such objective will therefore represent a milestone in the evolution of current satellite systems and certainly testify the entering the era of next generation satellite system. As a matter of fact, current systems are able to provide more limited capacity in the range of 10-200 Gbit/s (e.g., Hylas 2, Ka Sat, and Echostar XIX, just to cite a few) [52], whereby important re-engineering effort is needed to advance the design of satellite systems. In particular, (at least) three main technical challenges must be properly addressed in the overall system design[100, 75]:

1. Bandwidth scarcity;
2. Large gateway network;
3. Matching user needs in terms of provided data rate on the long term.

As to the first point (1), a suitable approach is to exploit the whole Ka frequency band (17.7-19.7 GHz) [59] for the user link and operate the feeder link in EHF frequency band (> 40 GHz) to make use of larger spectrum portions. In particular, some attention has been paid to the use of Q/V frequency bands possibly also combined with either higher frequencies (e.g., W frequency band) or optical feeder link to further increase the available capacity. On the other hand, it is worthwhile to recall that links operating in EHF or even more with free-space optics technology can be severely affected by atmospheric impairments. For instance, services such as teleconference and telemedicine, which are subject to high availability requirements, could be subject to important quality of experience degradation at users' level. Moreover, mission-critical applications (i.e., requiring small latency delivery and

no information loss such as alerting/warning and telemetry services) could suffer from important performance limitations, introduced by feeder link outage events, though of short duration. As to gateways operating with RF links, performance degradation occurs because of heavy channel fading introduced by rain, whereas links built on free-space optics technology are typically affected by signal blockage events due to the presence of clouds.

As a consequence, the design of the ground segment must be based on redundancy and space diversity concepts, hence additionally increasing the size of network [52], as also claimed in point (2). In particular, the need for large gateway network inherently comes from the target of achieving Terabit/s capacity, which requires a large number of gateways operating in RF. This aspect must be duly taken into account while targeting availability figures higher than 99.5%, whereby smart gateway diversity techniques (as further elaborated in the next section) are actually necessary.

The last point (3) is certainly one of the most important, i.e. matching the user needs, since it entails aspects of both ground and space segments. On the one hand, providing users with larger data rates than actually provisioned so far will require satellite networks organized with a much larger number of beams (> 200), which in turn will also result in narrower [75]. Moreover, in order to take advantage of the available frequency spectrum as much as possible, it is also envisioned to increase the frequency reuse factor, hence requiring proper interference mitigation techniques and scheduling solutions. On the other hand, meeting users' requirements in terms of data rate is tightly connected to the problem of resource allocation, actually consisting of minimizing the difference between offered and requested capacity. The solution of such a problem is actually to be implemented on both space and ground segments, in order for the satellite payload too to properly allocate resources (time/frequency and power, where meaningful). Design of proper radio resource management solutions is therefore necessary, as the current ones mostly based on static assignment are expected to perform poorly in presence of large satellite networks (i.e., > 200 beams) and high capacity availability.

As a result, the design of a broadband satellite network has to carefully take into account a number of key system parameters that have to be optimized in order to maximise throughput and availability figures (at least), just to cite a few fundamental performance targets. In more details, the following system parameters deserve a particular attention during the exercise of dimensioning a full satellite system and hence meet specific service requirements [75]:

- The user link bandwidth B , which is the amount of the Ka-band radio frequency bandwidth assigned to the user link and determined mainly by regulatory constraints;
- The number of beams N_b , which depends on the size of the coverage, the on-board user link antenna size and the beam pattern cross-over points;
- Number of colours n_c (directly correlated to the frequency reuse), i.e. the number of unique combinations of frequency sub-band and polarization. It determines the total system bandwidth $W(Hz) = N_b \cdot B / n_c$ and has an important impact on the level of co-channel interference;
- Number of on-board HPA's, which is limited by the on-board mass, power, accommodation and thermal dissipation constraints. Its value depends on the number of beams and the number of beams per HPA.
- Total bus DC power PDC: only a portion of the total DC power is allocated to the telecom payload (in particular to the Forward Link TWTA's) and converted to RF power.

2.1.2 Spectrum Regulations

As aforementioned, the amount of spectrum available for the HTS system has to be in line with the regulatory constraints as identified by the ITU worldwide, CEPT in Europe and also the different national authorities [52, 75, 100]. In the user link, due to the nature of the satellite signals from/to small user stations, coexistence with other services is very challenging and therefore regulations shall guarantee a certain

level of protection. In this sense, European decisions have been made to guarantee protection from interference and also exemption of individual terminal licensing in the following bands (that will be called “exclusive band” hereafter):

- 19.7 – 20.2 GHz for Space-to-Earth communications;
- 29.5 – 30 GHz for Earth-to-Space communications.

Additional decisions identify several other portions of spectrum that could be used for the uplink (Earth-to-Space communications) of the user link namely: 27.5 – 27.8285 GHz; 28.445 – 28.8365 GHz; 29.4525 – 29.5 GHz; 28.8365 – 28.9485 GHz.

For the downlink, the portion of 17.3–17.7 GHz could be considered as a candidate band under the umbrella of High-Density Fixed Satellite Service (HDFSS). However, the operation in this band shall be done without prejudice to the BSS feeder link service. In the following, these additional portions of spectrum will be referred to as “extended band”.

For the feeder link, two options exist:

- Ka-band: ITU Radio Regulations allocate the band 27.5 – 29.5 GHz on a worldwide basis for Fixed Satellite Services in the Earth to Space direction and the 17.7 – 19.7 GHz band in the Space to Earth direction;
- Q/V band: in the area from 37.5 to 43.5 GHz and from 47.2 to 51.4 GHz.

2.1.3 Smart Gateway Diversity (SGD) Architecture

The system design has to cope with feeder link impairments because of bad weather conditions that are particularly severe at high frequency bands. Adaptive Coding Modulation (ACM) alone may be unable to guarantee the required Quality of Service (QoS) system specifications, because of the very high signal degradation occurring in these bands (deep fading due to meteorological effects), which implies an undesirable throughput reduction versus feeder link channel propagation attenuation. Based on these facts, the Gateway Diversity (GD) principle was developed, which employs a set of (inter-)connected GWs via

terrestrial links. When the feeder link of a GW experiences a deep atmospheric fading, its traffic can be rerouted to another GW using the terrestrial network segment [59]. Diversity techniques exploit the spatial diversity existing among sufficiently-separated GWs in order to relax the margins in the link budget.

The multi-GW architecture achieves the desired availability at the expenses of redundant GWs added with respect to the N active ones. Each active GW is connected to a backup GW, which is activated if the main one undergoes unacceptable levels of fading. During this fading event, all traffic is rerouted to the redundant GWs; The rerouting decision will be taken by the Network Control Center (NCC). The additional GWs are deployed at least 20 km far away from the main ones in order to de-correlate the related (atmospheric) fading events. The user terminal has to lock its carrier to the redundant Gateways (GWs) after the diversity handover. If the system is designed to use N GWs, the mentioned basic diversity scheme requires $P \equiv N$ redundant GWs to achieve the desired level of availability; this scheme is denoted as the classical diversity approach. The main drawback of this classical GW diversity scheme is that we need to double the ground segment and the underlying costs and this could be considered unacceptable for future satellite systems. Therefore, it becomes important to adopt improved solutions to reduce the number of GWs needed. This is the reason why SGD (Smart Gateway Diversity) schemes have been proposed in [92, 58]. We refer below to the SGD techniques described in the patent in [10] to achieve an availability level similar to that of conventional GW diversity schemes, but requiring a lower number of backup GWs.

The SGD concept is based on the following basic aspects:

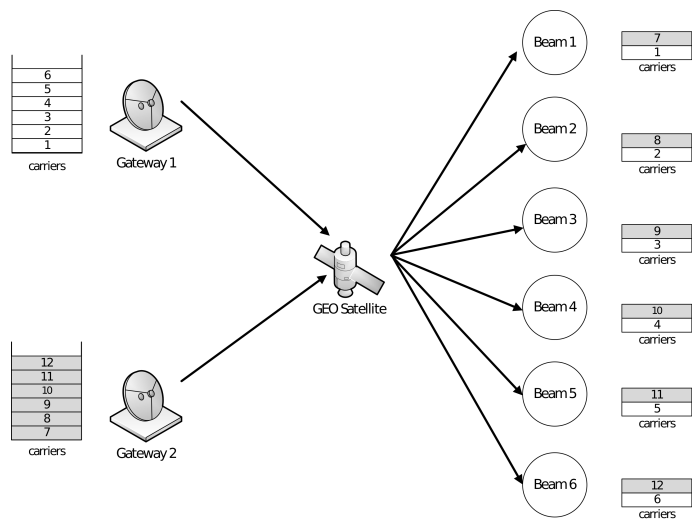
- All GWs are inter-connected by means of a terrestrial network;
- Each user beam can be serviced by a number of GWs (i.e., not only one, even if only one is actually used);
- In the event of a GW experiencing outage by a given user or reduced capacity, some or its entire traffic can be redirected on the ground towards another GW.

Different SGD schemes can be categorized as detailed below, depending on the use of P redundant GWs with respect to N system GWs [10]:

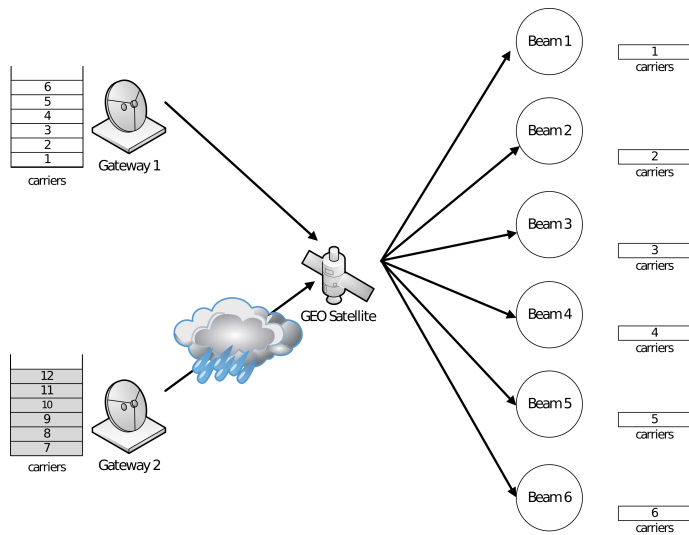
- Smart $N+0$ GWs,
- Smart $N+P$ GWs,
- Smart $++N+P$ GWs.

$N+0$ Site Diversity: The basic idea behind the $N+0$ macro-diversity scheme (depicted in Figure 2.1) is that each user beam is served by carriers coming from all the GWs or by a subset of GWs. This scheme requires a modification at the satellite payload level in order to filter (separate) carriers and recombine them from different feeder links. Each user beam is served by more than one GW through different carriers as shown in Figure 2.1a. A given user is typically locked to one carrier only. Impairments on one GW (see Figure 2.1b) will imply that all users served by the carrier will have to be migrated to the other carriers serving that beam and belonging to different GWs, which will not result in a full outage, but rather in a throughput reduction in the beams affected. The main advantage is in that no redundancy gateways are needed differently from the other proposed solutions (see below).

$N+P$ Site Diversity: This scheme, depicted in Figure 2.2, envisages a satellite payload with routing capabilities. Some extra GWs (redundancy) are needed to ensure that there will not be any reduction of capacity for users when a GW experiences an outage event. Unlike the previous scheme, this technique achieves a better throughput distribution (vs. fading) at the expenses of the need of P redundant GWs, as shown in Figure 2.2a. Basically, like in the ‘ $N+0$ ’ scheme each user is served by a number of gateways through different carriers. When one GW becomes unavailable, instead of losing part of the beam throughput, those carriers are switched at payload level. In practice, there are N concurrently-active GWs (at 100% of their carriers). Moreover, GWs are interconnected to facilitate the rerouting of traffic when a GW experiences fading conditions (Figure 2.2b) and there is the need to perform a GW handover, thus activating one of the P spare GWs:



(a) Clear-sky



(b) Rain.

Figure 2.1: N+0 Site Diversity

all the input traffic of the affected GW has to be rerouted towards a redundant GW. This SGD technique can manage the situation only if the number of affected GWs is lower than or equal to P . On board of the satellite (payload), there is the need of an $N+P:N$ non-blocking switching matrix.

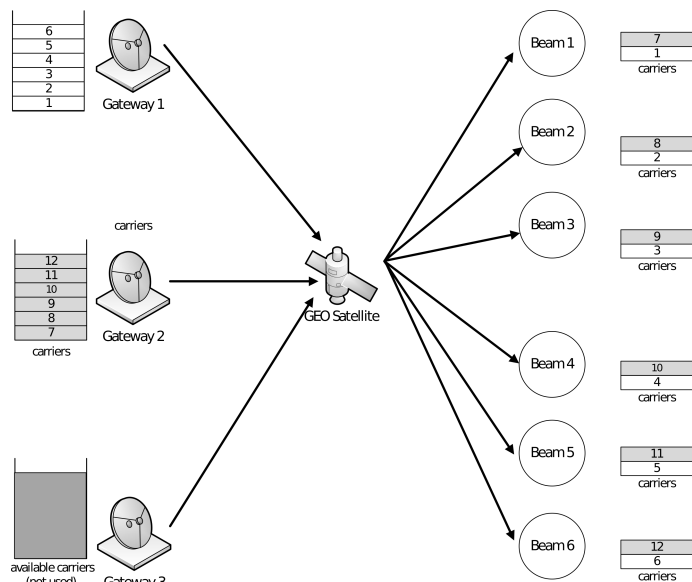
$++N+P$ Site Diversity: This scheme (depicted in Figure 2.3) is similar to the $N+P$ GW diversity, except for the fact that in nominal operating conditions all the GWs (e.g., all $N+P$) operate concurrently (Fig. 2.3a). Moreover, similarly to the previous two approaches every beam is served by more than one GW. This concept improves the system availability at the expenses of a more complex satellite payload. Let us consider an example with $N = 2$, $P = 1$, and 27 carriers (9 for each gateway). All $N+P = 3$ GWs work simultaneously using only $N = 6$ carriers (even if they can support up to 9 carriers). If one GW experiences outage (Figure 2.3b), its carriers are activated in other GWs (now each of the remaining 2 GWs is using 9 carriers) that are switched to the same user beams as before by means of a suitable reconfiguration of the on board switching matrix of the satellite. Thus, we can maintain the same capacity and connectivity as before. The satellite payload is more complex with this scheme.

2.1.4 Payload architectures

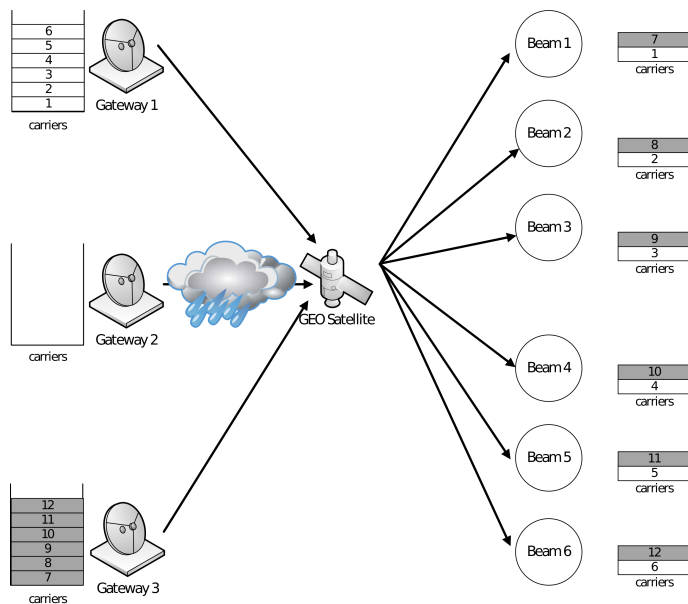
According to the type of user-link antenna and payload architectures, High Throughput Systems (HTS) can be divided into two main categories: 1) Single Feed Per Beam (SFPB) and 2) Multi Feed Per Beam (MFPB). The first category is currently the most popular, but the second category has had and will continue to have an important role for specific mission scenarios.

Single Feed per Beam (SFPB)

As depicted in Figure 2.4, in SFPB architectures, each spot beam on ground is generated by only a single antenna feed element. This means that the feed array has a number of elements equal to the number of beams. Given the conflicting requirements between beam size and

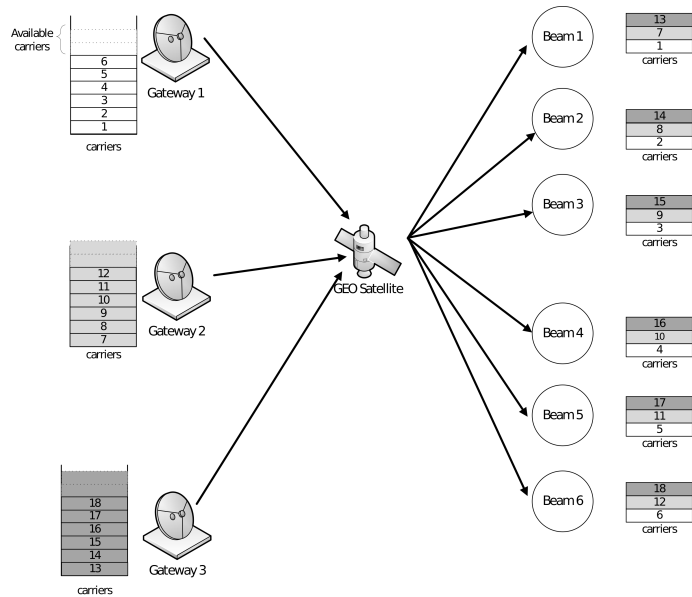


(a) Clear-sky

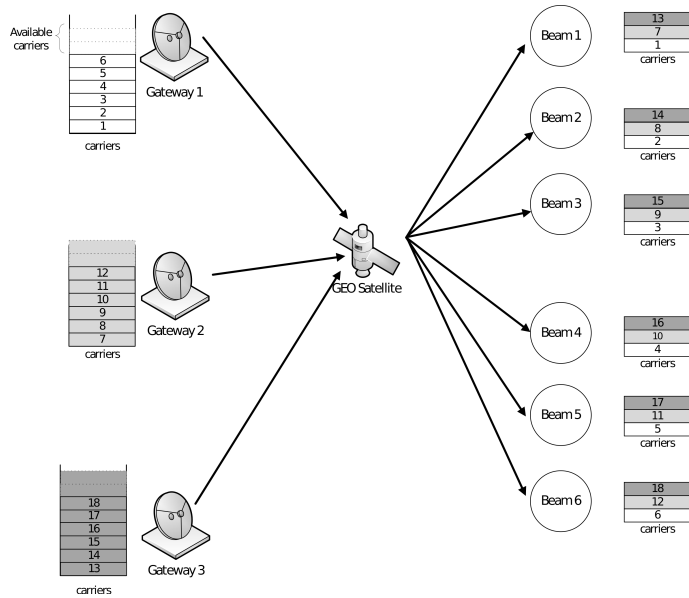


(b) Rain.

Figure 2.2: N+P Site Diversity



(a) Clear-sky



(b) Rain.

Figure 2.3: ++N+P Site Diversity

beam spacing in order to generate good antenna beam pattern, this solution typically foresees 3 or 4 separate antennas, each generating an interleaved set of beams within the coverage.

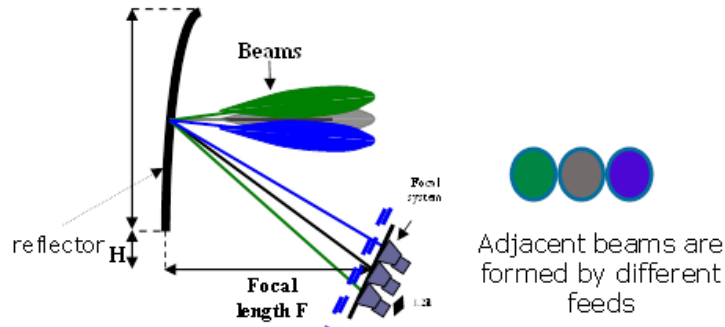


Figure 2.4: Single Feed per Beam (SFPB) Antenna concept.

The payload architecture associated to SFPB antenna is shown in Figures 2.5 and 2.6 for the Forward and Return link mission respectively. The payload architecture associated to an SFPB antenna can be described in terms of high level functional architecture without redundancy and is limited to the beams served by one GW. A possible assumption is that the coverage uses a four colour scheme, i.e. the user bandwidth (typically 500 MHz in Ka-band) is split into 4 sub-bands, each of these sub-band being regularly used within the coverage with a distribution aiming at limiting the system co-channel interference. In this particular example, one polarization is used for both forward and return user links. Other possible schemes foresee the usage of dual polarization whereby the user colouring scheme includes also the two polarizations, so that the user bandwidth is split into two sub-bands instead of four. Also, on the GW link, the schemes considered here assume to use one polarization, while, very often, two polarizations are used to exploit at most the available feeder link bandwidth. However, in both cases (user and feeder link), the extrapolation to the dual polarization equivalent is straightforward.

In the forward link mission, the carriers which can be uplinked by the GW are received on-board by the payload receiver section (filter

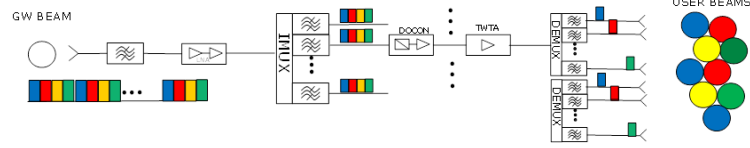


Figure 2.5: Example of Forward Link Payload Architecture for Single Feed per Beam (SFPB) Antennas.

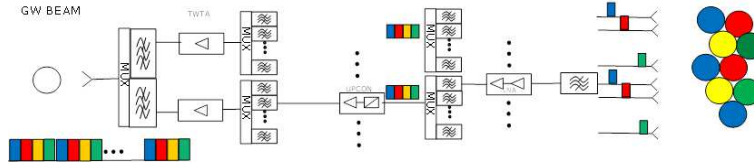


Figure 2.6: Example of Return Link Payload Architecture for Single Feed per Beam (SFPB) Antennas.

plus Low Noise Amplifier (LNA)) and then split into several groups of carriers, each serving an High Power Amplifier (HPA), which is typically a Travelling Wave Tube (TWTA) when Ku or Ka-band are considered. This function is carried out by the Input Multiplexer (IMUX) block and is required in order to feed each different HPA with a set of carriers which, after proper down-conversion, are translated to the same user link band. After power amplification, a Demultiplexer filter separates the different carriers (each representing one colour) so that each is fed to a different antenna feed element to create a unique beam. The opposite processing is carried out for the return link direction as shown in Figure 3. Here, the scheme of aggregation of carriers into the feeder link TWTAs depends on the feeder link budget and bandwidth. A different number of TWTAs may be required to properly amplify the total feeder link bandwidth.

An example of HTS implementing the SFPB architecture is the Eutelsat Ka-Sat [35]. This satellite has 4 SFPB reflectors each illuminating one fourth of the total 82 beams. The user bandwidth of 500 MHz is reused within the coverage with a four colouring scheme which also exploits double circular polarization so that the beam bandwidth is 250 MHz. The total delivered capacity of this HTS is around 90 Gbps.

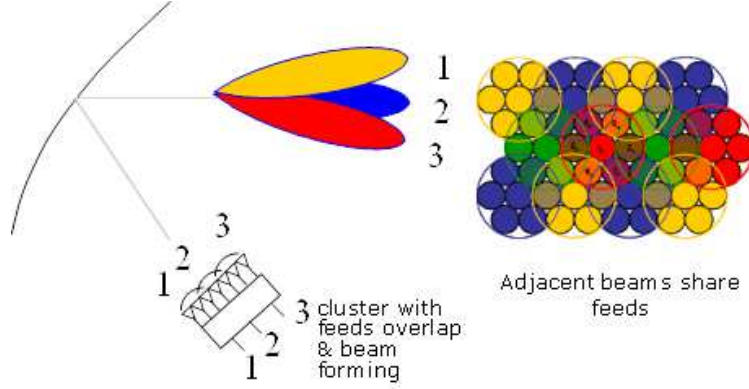


Figure 2.7: Multi-Feed per Beam (MFPB) Antenna concept.

The reader is referred to [35] for more details about this system.

Multi-feed per beam (MFPB)

In MFPB architectures, each beam within the coverage is generated by a number of radiating elements within the antenna's cluster of feeds. These number of feeds can be a subset of the cluster or it can be the whole cluster. In the example of Figure 2.7, 7 feeds are used to generate each beam. In this case, due to the limited number of elements generating a beam, a reflector is used to focus the antenna beam pattern. These architectures are also called **Focal Array Fed Reflector (FAFR)**, while, when all the feeds contribute to generate each beam (thus without a reflector) we name the architecture **Direct Radiating Array (DRA)**. Note also that, typically, some feeds contribute to the generation of multiple beams (two beams in the example of Figure 2.7).

Usually, with this kind of architectures, only one aperture is used to generate the whole coverage, thanks to a different trade-off between feed size and antenna beam properties. Therefore this architecture allows to save in terms of number of on-board antennas, while, however, the antenna itself and payload might result to be more complex due to the need to embark some sort of Beam Forming Network (BFN). This

is indeed required to excite the feed array with the signal amplitude and phase which result into the proper in-space combination of the different feed contributions that in turns create the wanted beam pattern. Thanks to the on-board BFN, these systems can often (depending on the specific architectures) implement a high degree of flexibility in payload resources and coverage. In the following, FAFR and DRA architectures are illustrated.

Focal Array Fed Reflector (FAFR). The FAFR architectures can be divided into two categories according to whether the BFN is implemented after or before the power amplification section. We then talk in terms of a *High level BFN (HL-BFN)* (*High level BFN*) or *Low level BFN (LL-BFN)* solution, respectively:

- The *HL-BFN* solution (see Figure 2.8 for a FW link mission) foresees a number of HPAs equal to the number of beams to be generated. As these are normally less than the number of feeds, these architectures have the advantage to embark a limited number of HPAs. However, the foreseen BFN is typically quite bulky with limited or no flexibility due to its high power nature. In addition, particular care has to be taken with its design as any power losses degrade the antenna Effective Isotropic Radiated Power (EIRP) and thus the overall system efficiency.

For this reason, the BFN is typically attached directly to the Feed Array so to avoid the losses due to the waveguides. An example of an implemented architecture is the AIRBUS MEDUSA [90] feed/BFN sub-system (Figure 2.9). This array can feed reflector of different sizes, according to the wanted coverage and beam sizes. The number of feeds used to generate each beam is seven, with re-use of feeds to generate adjacent beams. A typical MFPB payload employs two antennas, one for the forward and the other for the return user link.

The payload architecture associated to a HL-BFN FAFR antenna (Figures 2.10 and 2.11) is very similar to the one of the SFPB solution. The only difference is represented by the antenna where

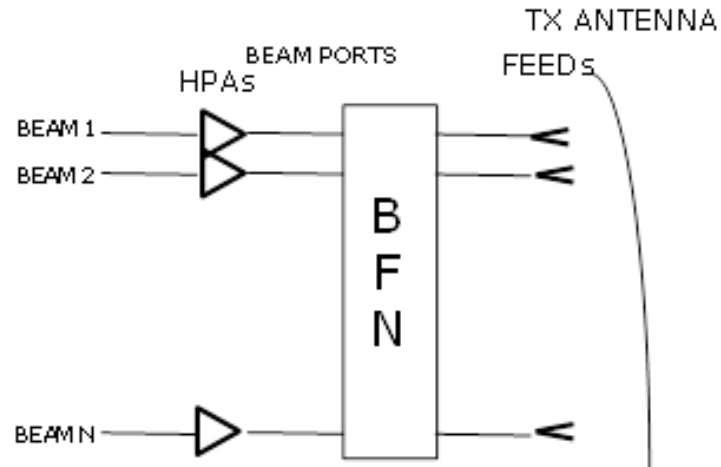


Figure 2.8: HL-BFN FAFR Transmit Antenna Concept.

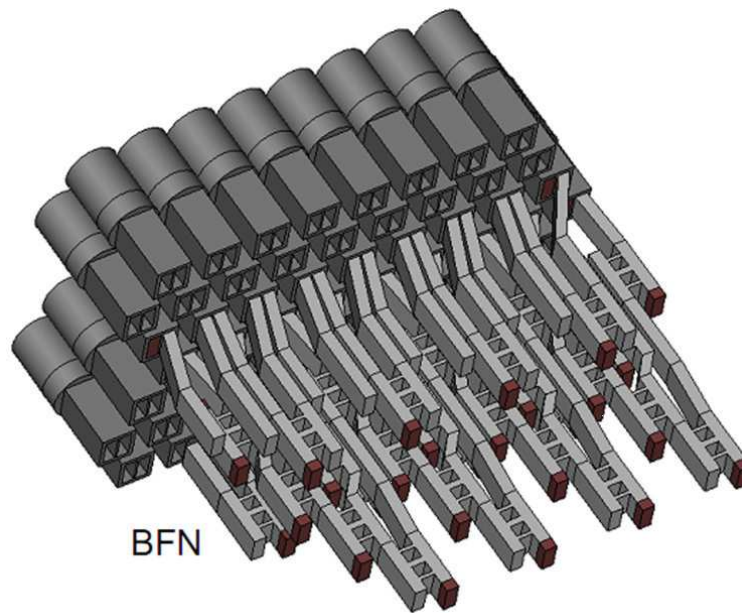


Figure 2.9: The layout of the MEDUSA Feed/Beam Forming Network.

(in the Forward link) instead of taking the DEMUX outputs di-

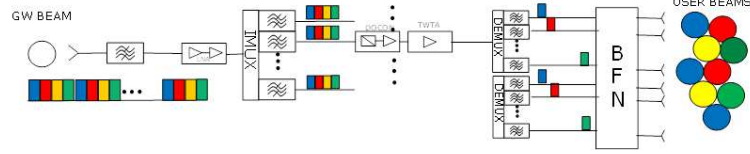


Figure 2.10: Example of Forward Link Payload Architecture for HL-BFN FAFR Antennas.

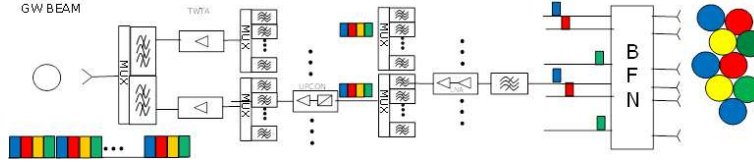


Figure 2.11: Example of Return Link Payload Architecture for HL-BFN FAFR Antennas.

rectly at feed level as in in the SFPB scheme, it foresees the BFN stage in between. Similar analogies apply for the return link direction.

- Within *LL-BFN* architectures (Figure 2.12), on the other hand, the BFN is implemented at low power. The number of HPAs is equal to the number of feeds, so that a relatively high number of them can be required particularly for large missions. Also, each HPA amplifies a composite signal made up of several carries and thus it has to work in relatively large back-off w.r.t saturation. However, the BFN, being low power, does not impact the antenna EIRP with its losses (the losses can be compensated for by a proper HPA pre-amplification stage) and can be made quite flexible (reconfigurable while in orbit) and possibly dynamic. In some implementations, the BFN can also be implemented digitally within an On-Board Processor (OBP).

An example of such architecture is the forward link mission of the Eutelsat Quantum satellite [97] (see Figure 2.13). This antenna is designed to provide up to 8 independent beams (4 in horizontal polarization and 4 in vertical polarization) and will in particular be capable of providing independently reconfigurable

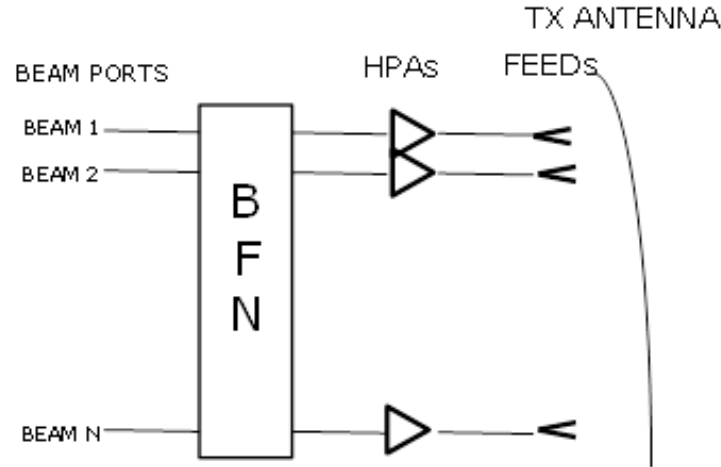


Figure 2.12: LL-BFN FAFR Transmit Antenna Concept.

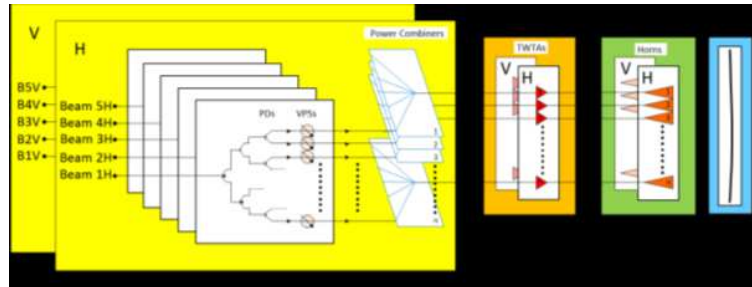


Figure 2.13: Quantum Transmit Antenna.

service areas. The downlink antenna provides full RF power sharing between the four beams in each polarization. Thanks to a low power BFN located before the TWTAs, the downlink antenna allows several beams to be generated without impacting the RF output section loss. The number of feeds has been baselined as 21 per polarization.

A possible high level payload architecture to be used with a LL-BFN FAFR transmit antenna described in Figure 2.14 (FW link) can be summarised as follows. The forward link carriers generated by one gateway are received on board by the classical re-

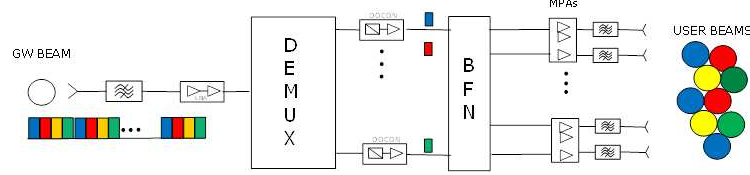


Figure 2.14: Example of Forward Link Payload Architecture for LL-BFN FAFR Antennas.

ceiver (filter plus LNA) unit. A demultiplexer is then used to separate the carriers which have to be routed to the different beams. This demultiplexer is typically implemented in several parallel units. A down-conversion stage then shifts the beam carriers to the right user link frequency, according to the desired frequency colouring scheme. The BFN synthesizes a linear combination of the signals fed to its input ports so that, together with the antenna, the right beam pattern is generated on ground. Following the BFN, the power amplification stage brings the feed signals to the right power level for the target antenna EIRP. This amplification stage can be implemented in both TWTA and Solid State Power Amplifier (SSPA) technology. For large systems, typically the SSPA solution is preferred due to the mass budget advantages. However, the power efficiency of SSPAs at Ku and Ka-band still represents an area of improvement, when compared to the one achievable by TWTAs. Therefore, this architecture requires careful management of mass and power budgets, particularly for large networks. The power amplification stage is often implemented by means of Multi-Port Amplifiers (MPAs). Indeed, in a basic architecture where each BFN output is associated with an amplifier and one feed, the HPA level must be tailored to the excitation law of each feed. This results in inflexible power to beam allocation. Furthermore, a failure of one amplifier would lead to very significant degradation of one or several beams. Instead, the distribution of the power amplification using an MPA (power pooling) gives flexible management of the feed power levels and is more robust considering the failure of one HPA within each MPA.

Part of this architecture can be implemented in digital within an OBP. In particular, the DUMUX and the BFN units lend themselves well to a digital implementation, provided the computational complexity can be afforded.

Direct Radiating Array (DRA). Within a DRA solution no reflector is present, i.e. the antenna itself is represented by the feed array. Each feed is used to generate the complete set of beams within the coverage. Therefore the required BFN is typically more complex than the FAFR solution due to the beam ports to feed port full interconnectivity matrix and the high number of feeds. However, some solutions exist to minimize complexity. For example, various schemes have been developed to minimize the number of array elements given a certain aperture size as well as the complexity of the BFN [11, 21, 64].

A rather famous example of DRA solution is represented by the Spaceway antenna where a DRA could generate 24 beams (12 LHCP + 12 RHCP) with a fully-agile digitally-controlled analog BFN. The number of radiating elements was 1500 and thus 3000 the SSPAs and used square-aperture high-efficiency dual-polarization horn radiators.

Another interesting example is the receive DRA which will be embarked on the Eutelsat Quantum satellite [97] (Figure 2.15). This Ku-band dual polarized antenna has 84 elements generating 4 independently flexible coverages on each of X and Y polarisations, integrated LNAs, attenuators and phase shifters, filtering, (PSU, ICU, heat control, redundancy).

2.1.5 Flexible HTS Payload Architectures

The capability to flexibly allocating the satellite payload resources over the service coverage is becoming a necessity for next generation broadband satellites employing a high number of spot beams. Indeed, previous and current broadband systems have shown that large multi-beam HTS are typically able to fill-up fairly quickly the capacity of some beams, while some others remain (almost) empty over a relatively long part of the satellite life time. The consequence is a loss of satellite operator's revenue due to the number of customers lost within the hot-spots

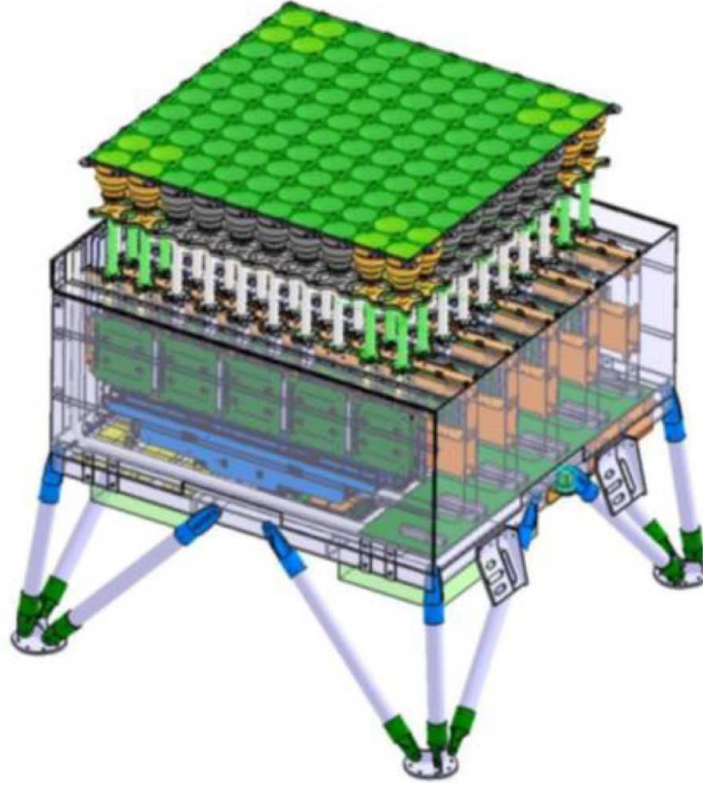


Figure 2.15: Eutelsat Quantum satellite uplink DRA.

(filled-up beams) and the waste of resources over the empty spots. A number of initiatives have been recently taken by the satellite industry in the direction of missions supporting flexibility at payload equipment level. The Hylas satellite [97] and the more recent Quantum [12], are just a two examples.

In order to correctly characterizing the performance of a broadband HTS system the following definitions are in order:

- The capacity demand is the capacity that is requested by the users which is typically geographically non-uniform and time variant.
- The offered system capacity represents the maximum capacity

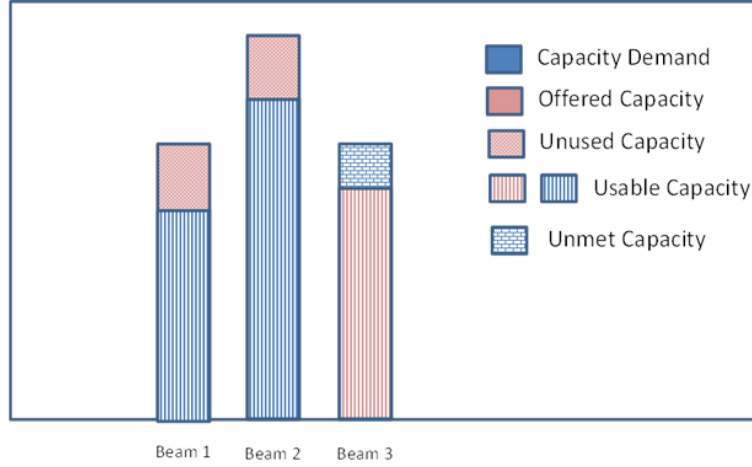


Figure 2.16: Definition of the different network capacity values.

of the system, while considering an infinite capacity demand per location.

- The usable system capacity is the capacity that is really sold taking into account the real capacity demand per location.
- The unused system capacity is the difference between the offered capacity and the usable system capacity.
- The unmet capacity demand is the difference between the capacity demand and the offered capacity.

These definitions are pictorially explained in the Figure 2.16 where an example with a hypothetical system with three beams is represented.

The primary goal of flexibility is to minimize the unused and unmet capacity. The introduction of flexibility helps a satellite operator to manage the risks accounted by the unpredicted changes, like regulatory, competing, and socio-economic contexts. Flexibility refers to the ability to change the configuration of the system during the operational life of the satellite. In the following, we will focus on the flexibility in the forward link of HTS systems, as it clearly represents the most important link in determining the revenues of the operator.

A number of techniques are available to support flexibility. In addition to flexibility in coverage (which will not be explicitly addressed here), the following is the list of such techniques: i) beam hopping, ii) flexible bandwidth allocation, and iii) flexible power allocation.

Beam Hopping

This technique is exactly dual w.r.t the flexible bandwidth allocation technique, i.e. it can be explained by replacing time with frequency. Indeed, Beam Hopping (BH) foresees that different co-channel beams served by the same HPA get assigned different time slots. By modulating the duration of the timeslots, different offered capacity values can be reached in different beams. This is schematically depicted in Figure 2.17, which shows a possible architecture for implementing BH within a SFPB payload architecture implementing 8 beams, 4 in RHCP and 4 in LHCP. The high power switches at the HPA output commute from one throw to the other synchronously with the carrier time slot framing, so to route the right time slot to the right antenna feed and ultimately the right beam. It is clear that a network synchronization mechanism has to be in place between the gateway and payload so that the on-board switches can correctly commute between two consecutive time slots. The system co-channel interference is kept low by making sure that the nearest beams are either illuminated in different time instants or served by the orthogonal polarization.

An important parameter of the BH technique is the Switch Throw-count S_c , which defines the number of positions each of the switch poles can be connected to. Together with the number of HPAs (N_{HPA}) and number of beams (N_b), represent the key system parameters that define the performance of the satellite network. Let's us now consider Figure 2.18 which depicts an example of the so called system illumination table under the assumption of regular (left-hand picture) and irregular (right-hand picture) illumination). The illumination table represents on the horizontal axis the fraction of the illumination period (TL) associated to a given beam, while on the vertical axis the number of beam simultaneously illuminated is represented by means of a number of rectangles. Therefore, the rectangles on the same row represent

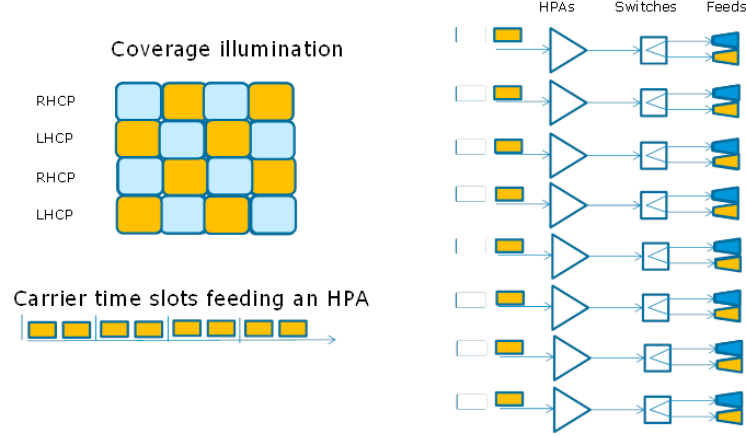


Figure 2.17: Illustration of the BH concept with a SFPB Architecture (transmit direction).

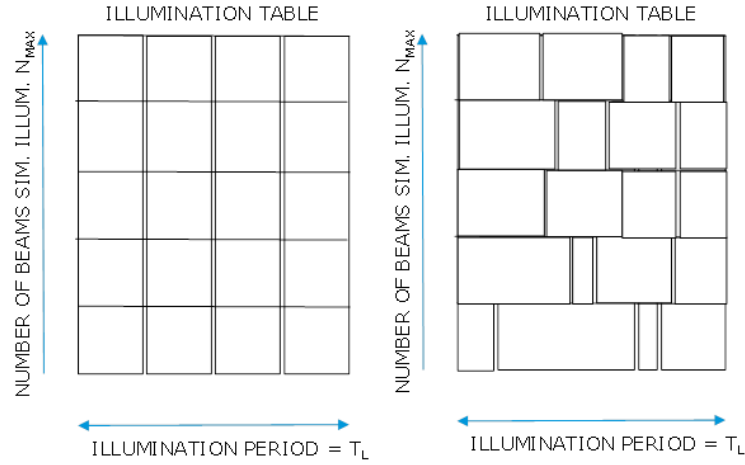


Figure 2.18: Example of Regular and Irregular BH Illumination Table.

the beams that are illuminated by the same HPA, while the number of rectangles in a column represents the total number of HPAs.

It follows that the total number of beams can be written as $N_b = S_c \dot{N}_{HPA}$. This equation gives the number of beams that can be supported given a certain Switch Throw-count and the total number of on-

board HPAs. The irregular version of the illumination table describes the beam time slot allocation in case of variable traffic request per beam. To note that if no particular system countermeasures are taken, excessive co-channel interference may be caused particularly when in the presence of a cluster of hot spots (high capacity requests). In general, the duration of the time slot which can be allocated to a given beam is a multiple of a given time quantum which can be as small as the DVB-S2x Superframe (Annex E of [31]). Such framing structure has been devised in order to properly support the BH technique. Its general scheme is illustrated in Figure 2.19, which shows a number of consecutive time slots interleaved with guard time of dummy symbols, within a certain time period which repeats in time. This time period is the BH Illumination period which is typically set to be around tens of ms. The particular value is set as a result of a trade-off between the system latency/latency jitter (indeed a given beam gets traffic only once during an illumination period) and the resource flexibility requirement (the higher the number of time slots within the illumination table, the higher the potential beam traffic unbalance). Each time slot has to allow the receiver to demodulate/detect its content without the need to exploit any knowledge from previous time slots (as these might address other beams). Therefore the receiver channel estimation function, which aids the data demodulation and detection has to have the necessary means within the time slot to derive accurate channel estimates (amplitude, timing, carrier phase and frequency, SNR) within the duration of the time slot. This is done in the Annex E of [31] by proper apportionment of header and pilot symbols as described in Figure 2.20. More details about the superframe and its applications can be found in [31] and [32].

The guard time between time slots is set long enough to allow the payload switching to illuminate a different set of beams. This function is performed in different ways according to the payload architecture. In LL-BFN FAFR/DRA architectures, a possible approach is to change the BFN weight coefficients (amplitude and phase) as in [77]. The BFN complex weight coefficients $w_{i,j}$ can be controlled digitally via an on-board Beam Weight Processor which computes and load them to the BFN following the BH illumination table.

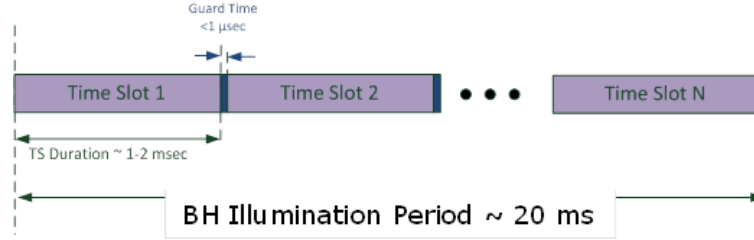


Figure 2.19: Example of framing required by BH.

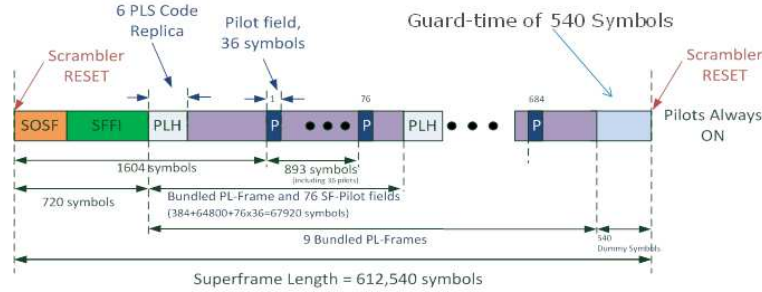


Figure 2.20: DVB-S2x Annex E Superframe format supporting Beam Hopping.

In other architectures (SFPB for example in Figures 2.21 and 2.22) the switching is carried out by a high power switch unit as the ferrite switch. This unit, which typically weighs a few tens of grams, is controlled by a Ferrite Switch Control Unit which is in charge of commanding the commuting of the payload switches according to the desired BH illumination table.

A further elaboration of the BH concept is provided in [76] (Figure 2.23). This concept exploits BH to extend the flexibility also to the share between the forward and the return link throughputs. To this end, the payload is composed by a number of so-called “Pathways” which are the equivalent of a classical transponder but with flexible assignment to either the forward and the return link. This flexibility is supported by the RX and TX switches which, synchronously with the hopping table, connect at different time instants either the user beam feeds or the GW user beam feeds to a given pathway. The result is that the pathways are partially shared between the forward and return links.

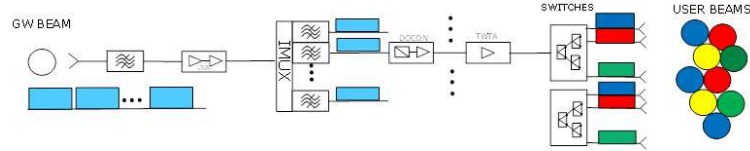


Figure 2.21: Example of Forward Link SFPB Payload Architecture supporting BH.

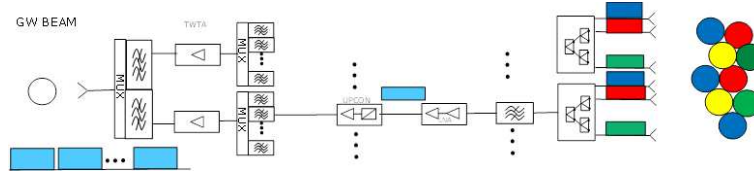


Figure 2.22: Example of Return Link SFPB Payload Architecture supporting BH.

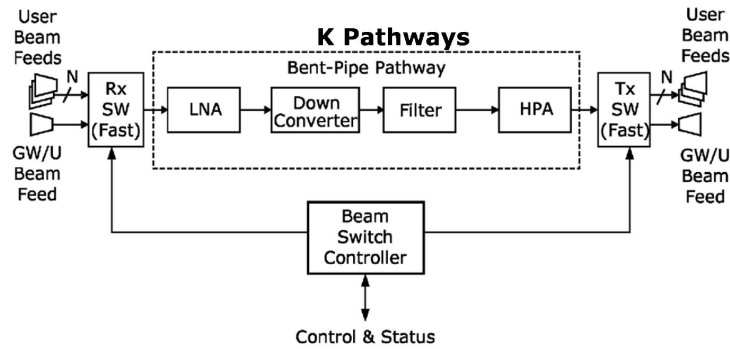


Figure 2.23: Flexible Forward and Return Capacity Allocation Scheme based on BH [76].

The fraction of time spent by the switch in each position determines the capacity provided to each beam as well as the capacity provided to the FW versus RT link directions. In order to support this scheme, the gateway and the user have to share the same bandwidth (which could be the whole Ka-band allocation for HTS) as well as the same antenna beam pattern.

Flexible Bandwidth Allocation

This technique consists in tuning the amount of band that is allocated to a given beam according to the relative capacity demand. In so doing, part of the bandwidth that is allocated to low demanding beams is transferred to high demanding beams. This scheme is sometimes called “Irregular Frequency Re-use” as the band allocated to each beam is variable. As a consequence, similarly to BH, if no particular system countermeasure is taken, excessive co-channel interference may be caused particularly when in the presence of a cluster of hot spots (high capacity requests). This technique can be implemented following the three approaches outlined in following, i.e. by using either a *High Level/Low Level Switch Matrix* or by means of an *active antenna* (DRA or LL-BFNFAFR):

- *High Level Switch Matrix.* Irregular frequency re-use can be achieved by splitting unevenly the user bandwidth allocated to the two beams served by the same TWTA (which is a typical configuration for a four color scheme network) and flexibly routing the two portions of the bandwidths to different antenna feeds (see Figure 2.24). The drawback of this approach is that the switch matrix implies EIRP losses as well as a relevant impact to the payload mass budget. Overall, this reduces the overall satellite flat throughput. In addition, since the granularity in throughput flexibility is determined by the bandwidth of the switched carriers, when choosing a narrow-band carrier, the complexity of the switch matrix could grow too high, while viceversa the resource flexibility resolution might not be meeting the target requirement.
- *Low Level Switch Matrix.* The approach using a Low Level Switch matrix is properly exemplified by the Airbus product called Generic Flexible Payload (GFP) [85], which comprises a number of key blocks. Agile Integrated Downconverter Assembly (AIDA) provides flexible down-conversion from the uplink beam frequency to a common intermediate frequency of C band. The second equipment in the GFP is Routing and Switching

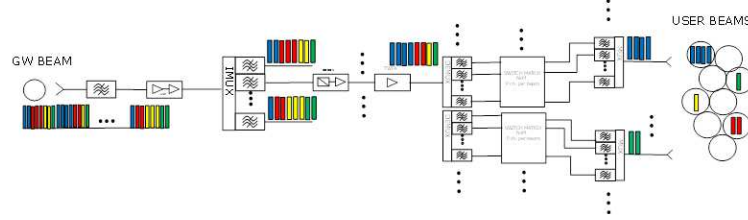


Figure 2.24: Example of Forward Link SFPB Payload Architecture supporting flexible bandwidth allocation with a HL-Switch matrix.

Equipment (RASE) (the low level switch matrix). The RASE operates in the C band uplink frequency and provides flexible connectivity between uplink and downlink beams. This combination of modules allows any uplink beam to be connected to any downlink beam. The last equipment is the Single Channel Agile Converter Equipment (SCACE). SCACE is an analogue processor which provides fully flexible channelization and channel amplification. SCACE offers the ability to control on a channel by channel basis, uplink centre frequency, downlink centre frequency (and hence frequency translation), channel bandwidth and channel gain including full channel amplifier functionality with FGM and ALC modes.

- *Flexibility with Active Antennas.* The most flexible solution, at least in principle consists in the use of a DRA or LL-BFN FAFR architectures to flexibly allocate a different bandwidth to different beams (Figure 2.25). To this end, the BFN beam ports are extended in number in order to take as inputs the carrier representing the flexibility granularity quantum. If more than one carrier quantum (say N carriers) is to be allocated to a given beam, the BFN is set to synthesize the same beam for the relevant N beam ports. Therefore, the granularity flexibility is a key factor determining the complexity of both the DEMUX and BFN. However, the output section is not affected by the choice of the flexibility granularity, including the EIRP performance.

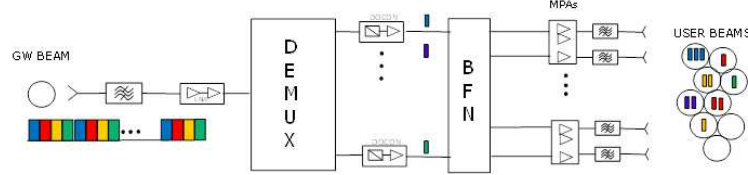


Figure 2.25: Example of Forward Link Payload Architecture supporting flexible bandwidth allocation with an active antenna.

Power Flexibility

To better match the capacity demand in each beam, a possible approach is to distribute the total amount of payload power unevenly across the different beams. Lower power would be assigned to beams with lower capacity demand, while higher power would be given to hot spots.

This technique can be implemented by means of the FlexTWTA technology where the saturated power of a TWTA is adjusted according to the capacity demand of beams served by the amplified carriers. Typically, the anode voltage is varied (while collector voltages are held constant), leading to a different cathode current and resulting in a different TWT gain and output power [51]. The result is that the power conversion efficiency of these devices remains almost constant within a range of $2 - 3\text{dB}$ of output power variation as opposed to a classical TWTA where the efficiency drops quite fast as function of the back-off. In case of one HPA serving two beams (which is a typical configuration), this technique works if the two beams have similar capacity demand. Alternatively, if the two beams have different capacity demand the power transfer from one beam to the other is done by suppressing part or all the carriers serving the low demanding beam.

Another approach for realizing flexible power allocation foresees the exploitation of MPAs. An MPA [71] is composed of an array of n power amplifiers in parallel and a pair of complementary $n \times n$ Butler matrix networks that consist of 900 hybrid networks [89]. A 4×4 MPA is shown in Figure 2.26. The signal at each input in the MPA is divided in n signals with particular phase relationships. These signals are separately amplified in each amplifier and are recombined in the output Butler

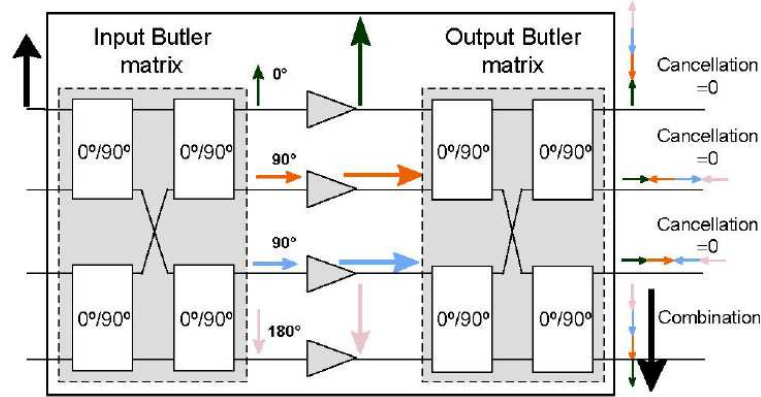


Figure 2.26: 4x4 MPA architecture, taken from [71].

matrix. The main advantage of this amplification architecture is that it provides intrinsic power flexibility, since power is shared between the channels. The combined power of all the power amplifiers is available for any channel provided that the other channels don't require power at the same time. This power flexibility is obtained without increasing the power consumption if an ideal MPA operation is considered.

The drawbacks of flexible power allocation are that any power variation has intrinsically a limited impact to the offered beam capacity due to both the inherent diminishing return behaviour of the Shannon capacity function (spectral efficiency versus power) as well as the presence of residual intra-system co-channel interference. Therefore power flexibility is often combined with either time or bandwidth flexibility. For example, a forward link payload architecture implementing BH as well as power flexibility by means of MPAs, can be considered and is described in Figure 2.27. The presence of MPAs allows the payload to support some throughput flexibility between the cluster of beams stemming from the same MPA, thus extending the flexibility contour of the payload.

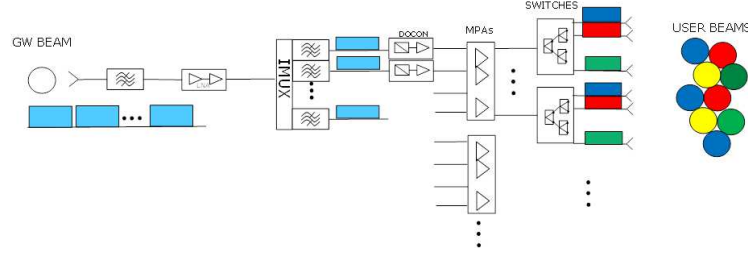


Figure 2.27: Example of Forward Link SFPB Payload Architecture supporting BH and MPAs.

2.1.6 Research Challenges

Radio Resource Management

The resource allocation problem in satellite networks has been always central in research and standardisation activities in the last twenty years. In spite of the large effort dedicated so far to this activity, no closed solutions have been found, unless model simplification were introduced, because of the large number of variables that are in play. As such, the resource allocation problem has been often reconsidered in the light of more classical problems such as scheduling, stochastic knapsack, and interference mitigation, just to cite a few. Moreover historically, resource allocation has been targeted from a ground segment perspective, like in [20] where the problem is addressed from a scheduling viewpoint, allocating different ModCods to the satellite beams.

On the other hand, more recently the resource allocation problem has started to be studied also from a space segment perspective, also owing to the recent advances in the satellite payload design, as illustrated in [9]. In more words, more sophisticated payloads have been developed in the last years[66], so as to cope with the time and geographic variations of the capacity requested by each beam. As a result of this system adaptation, three main payload operation concepts have been devised: beam-hopping, frequency-flexible, and power-flexible. As already pointed out in the previous subsections, the first makes use of a time-slotted illumination window so that it is possible to define the sequence of beam illumination and the number of slots assigned to

each beam according to the traffic demands and the antenna radiation pattern. The second is actually dual to beam-hopping as it provides flexibility in frequency, i.e. it is able to allocate bandwidth to beams in relation to offered and requested traffic. The third instead operates in the power domain, allowing to flexibly allocate power to beams as function of the conveyed traffic. It is also important to notice that the design concepts of power- and frequency-flexibility can be merged together in order to develop a more sophisticated payload. Likewise, the beam hopping concept can be also conjugated with power-flexibility.

Finally, it has to be remarked that all the three options have obviously important implications on the specific payload design (e.g., number of TWTA and structure of the payload connection matrix) and the related constraints (e.g., mass and available power) imposed by the technology available nowadays. In [61, 62] the problem of time/beam allocation is addressed and a closed form solution for the optimal resource allocation is derived in a simplified setup where no interference contribution is assumed. Moreover, two different utility functions are taken into account, in order to see how to 1) match the requested bitrate and 2) maximize the product of the ratios between the offered and requested capacity across the beams.

Still, [62, 23] elaborate the advantages of multi-beam with respect to single beam satellite systems with respect to different performance metrics. In particular, optimal power allocation is derived for two different objective functions, one leading to throughput maximization and the other related to fairness. These studies, however, neglect interference contribution leading to a partial representation of a real system. Such an assumption is relaxed in [24], where a phased array antenna is assumed at the satellite and call-admission control schemes are investigated. On the other hand, the advantages that may come from power allocation are explored in [87, 13], which consider a two-stage sub-optimal algorithm applied to solve a non-convex optimization problem, helping to shed some light in the relations between power allocation and offered traffic. Comparison of beam-hopping and flexible systems is provided in [33], where implementing non-uniform bandwidth allocation and making use of sizable beams is investigated to further improve the

performance of flexible payload systems. Finally, a more recent study about the potentials of bandwidth and power allocation in flexible payloads is provided in [26], where a resource allocation algorithm based on simulated annealing is developed to optimise an objective function properly tailored to the characteristics of the system's and users' QoS.

Open Issues

Open research issues concern the optimised control of the satellite payload functions, when interaction with service providers and satellite operators is needed to closely meet users' demands. From this standpoint, the current practice is actually based on proprietary solutions or limited support for the overall QoS management. In particular, the availability of a standardised protocol interface would allow satellite operators and service providers to use the satellite capacity more efficiently by seeking to match the users' requests. Apart from the technical complexity of such a problem (an optimal solution of the resource allocation problem is hard to find in a closed analytical form because of the non-convex shape of the objective function (see also [26])), the use of a standardised interface will be very helpful for the design of payloads of next-generation satellite systems. To overcome this possible performance limitation, the Satellite Communication and Navigation (SCN) working group within ETSI SES (Satellite Earth Stations & Systems) area has been working towards the preparation of a technical report able to collect the requirements from different satellite operators and elaborate on the main functionalities to be offered by a standardised interface, whose specification is however left to further discussion and activities within the working group.

2.2 Non-GEO Satellite Systems

2.2.1 Constellations

Traditional LEO Systems

Satellite constellations are not new; they have been proposed at the end of 90's as a possible alternative to GEO satellite platform to reach

the same level of coverage, though at much limited latency, hence possibly resulting a valid competitor to terrestrial networks. Nevertheless history tells that despite the initial attractiveness of this solution, the success has been limited to voice applications (e.g., Iridium [70]) and very few deployments became eventually stably operative because of business issues and the increasing service offers provided by 3G technology starting from 2005. In order to improve the robustness toward handover events and possible congestion events, some constellations also implement inter-satellite links that may form an aerial network with extended (broadband) coverage. Given the increasing appeal of free-space optic technology, it is envisioned that future constellations may have ISLs built on free-space optic technology, in order to take advantage of the very large data rate there provided.

Constellations can instantiate discrete (Internet) Autonomous Systems (AS), with individual internal networking operation, or integrated satellite-terrestrial networks where satellite and ground links are part of the same domain. In both cases, we notice two types of satellite links: inter-constellation and intra-constellation links. Inter-constellation links connect ground stations with the constellation, thus allowing the communication between terrestrial and satellite networks. Intra-constellation links connect two satellites within the same constellation, thus supporting the constellation network and providing an alternative (to the terrestrial ones) global wireless path. All in all, the main goal of constellations is to increase the coverage of satellite communications thus complementing ground links with various advantages, including relaxing the traffic load of terrestrial networks or in case providing the only possible network access option when the terrestrial infrastructure is missing (e.g., underserved areas, ships, aircraft, etc.) or temporarily not available (e.g., during disaster relief operations).

A constellation can rely on GEO, MEO and LEO satellites, which deliver different advantages based on of their altitude and orbit. The flight altitude is a key factor for the coverage range, since higher orbits permit observation of wider swaths, and to the energy requirements of the antennas, as longer communication links require more power to combat attenuation. Additionally, the orbiting plane affects

the satellite mobility with regard to Earth stations and, consequently, the frequency of intra-constellation handovers and the complexity of routing. For instance, GEO satellites provide large coverage and can effectively reduce the number of satellites composing a constellation of global range; three satellites can provide global coverage (though excluding polar regions). Additionally, GEO satellites are geostationary, thus avoiding mobility issues and simplifying the operation of inter-satellite handovers and routing decisions. However, due to their high altitude, GEO constellations also induce significant propagation delay, which can have a certain impact on delay-sensitive applications and can present challenges to the performance of transport protocols, such as Transmission Control Protocol (TCP).

On the other hand, LEO or MEO constellations consist of a set of satellites orbiting the Earth with high constant speed at a relatively low altitude. Therefore, the major advantages of LEO and MEO satellites are the reduced communication delay and lower power requirements for Earth-satellite communication. But the major challenge of LEO and some MEO (e.g., O3B) satellites is the increased complexity of the network protocols that need to address the mobile network of aerial nodes. Complex methods that deal with the orbital planes, the inter-satellite handovers, and the routing policies are required in order to provide continuous, fast and global communication. In the remaining of this section, we mostly discuss MEO and LEO constellations. Last but not the least, terminals should also have a tracking antenna, hence making the design and the implementation of the entire system with all its players significantly complex.

Constellation types Most satellite constellations are built upon the concept of streets of coverage. A street of coverage (Figure 2.28) is a set of satellites co-rotating in the same plane handing over communication to the consequent satellites. Multiple streets of coverage are exploited in order to enhance the coverage of the satellite constellation, thus providing a continuous and wide range service.

Streets of coverage are the base of the star constellation, a simple yet effective satellite constellation design that provides global cover-

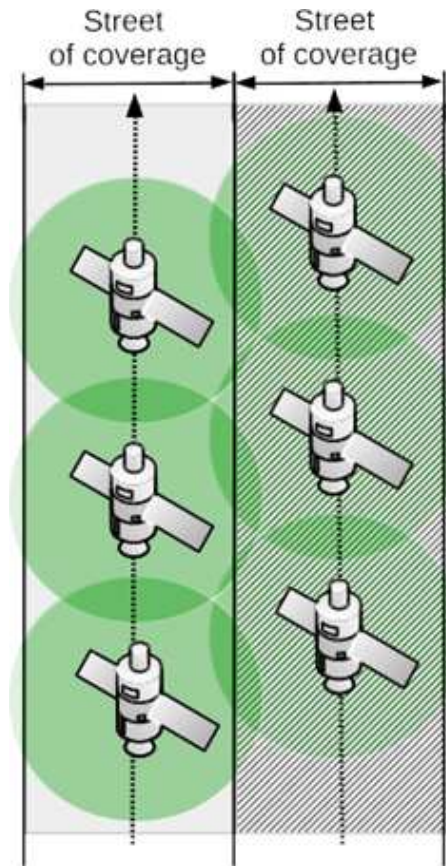


Figure 2.28: Streets of coverage.

age. In a star constellation (Figure 2.29), orbital planes are inclined at a constant angle, roughly vertical to the Equator, which is considered the plane of reference. Orbits overlap in the poles, forming a star that named the design, and achieving maximum dispersion near the Equator. Thereupon, the Earth is separated into two hemispheres with regard to the direction of the planes; in one hemisphere satellites travel towards the North Pole, while in the other they travel away from the North Pole. The per-hemisphere homogeneity of orbits is both a gift and a curse as it facilitates communication among adjacent streets of coverage but also challenges inter-plane communication in the “gray”

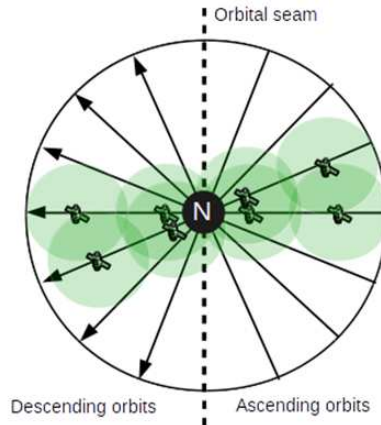


Figure 2.29: Star Constellation as seen from the North Pole. The arrows depict the satellite orbits, i.e., the streets of coverage. Only two streets of coverage are depicted.

area among two adjacent orbits of opposite directions. The Doppler shift is significantly increased throughout the seam and challenges the operation of intra-orbit links, thus fragmenting the global network. The network division complicates path selection in star constellations and highly affects communication latency. A second issue of star constellations is coverage variance across the prime meridian. As noticed earlier, streets of coverage overlap on the poles and reach maximum separation at the Equator, thus placing multiple satellites close to the former and a minimum satellite density at the later. Multiple satellite coverage penalizes resource utilization, but does not degrade the performance of the system. On the other hand, single satellite coverage of areas makes transmission sensitive to obstacles, such as tall buildings, therefore degrading the quality of communication. Considering that the majority of the human population, hence the demand for communication, is located at the mid-latitudes of the Earth, the coverage distribution in star constellations is a critical design parameter.

In order to address the demand-coverage mismatch a second type of constellation, called delta constellation (Figure 2.30), was introduced. In delta constellations streets of coverage are inclined at a constant

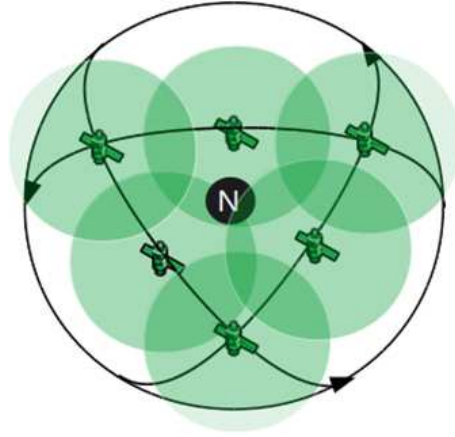


Figure 2.30: Delta Constellation as seen from the North Pole. The arrows depict the satellite orbits/streets of coverage.

angle, noticeably less than 90° at the Equator, which is considered the plane of reference, thus providing minimal, if any, coverage at the poles but maximum diversity near the Equator. This formation is named after the Greek letter Δ , because when the constellation employs a minimum of three orbital planes, a rounded triangle, similar to Δ , is formed around the pole by the planes. As a result, orbits towards and away from the North pole are overlapping, instead of being alienated into different hemispheres, thus providing increased and continuous coverage without cross-seam implications. Nevertheless, delta constellations also exhibit significant weaknesses, such as challenging inter-satellite communication. Even though service continuity is increased across mid-latitudes, the satellites traversing inclined orbits experience higher relative velocities and Doppler shift in comparison with inter-plane links for polar orbits, thus disallowing inter-plane networking. Expectedly, the satellites operate mostly as relays offering service within the boundaries of their coverage, thus limiting the relaxation of traffic load in terrestrial networks.

Inter-Satellite Links A satellite constellation can be studied as a typical network graph where vertices, or satellites, are joined by edges,

or Inter-Satellite Linkss (ISLs). ISLs may exploit space-free optical technologies for establishing communication, such as laser, which offer great throughput but also require line-of-sight propagation and precise transmitter-receiver adjustment. Based on these requirements, we distinguish three categories of ISLs: intra-plane, inter-plane and inter-seam. Intra-plane is the easiest to deploy type of ISLs as they connect satellites within the same orbital plane, where nodes maintain almost stationary relative positions. Therefore, the main requirement in order to establish intra-plane ISLs is to employ enough satellites per street of coverage in order to achieve line-of-sight contact. Inter-plane ISLs connect satellites of adjacent orbits that follow the same direction (supported only by star constellations). This kind of ISLs is more challenging to deploy than intra-plane since the distance between the adjacent planes oscillates throughout the prime meridian. Additionally, depending on the flying altitude, different number of orbits must be employed so as to guarantee continuous line-of-sight connection. Finally, inter-seam ISLs are the most challenging satellite links, trying to connect adjacent orbits over the seam of star constellations. The relative speeds of the satellites require frequent handovers and the increased Doppler often disallows the successful establishment of links, thus complicating routing and destabilizing the performance of constellations. The use of ISLs allows the creation of an alternative global communication (and content delivery) network that reduces the traffic load in terrestrial networks, mitigates the cost for terrestrial infrastructure and decreases the spectrum demand for satellite-terrestrial communication. Despite the increased propagation latency, the constellation network can offer high throughput dissemination paths that are suitable for bulk transfers such as video streaming, thus relieving the bandwidth requirements of the terrestrial network. In addition, intra-constellation routing relaxes the “1 Gateway and Terminal per satellite” requirement, as ground stations may be shared by all constellation satellites, thus reducing the deployment and operating costs of ground infrastructure. Finally, the reduction of ground stations also decreases the demand for radio spectrum for Earth-satellite communication, thus avoiding interference problems that arise when multiple systems utilize the same spectrum.

However, even in the absence of ISLs, constellations still present some advantages in terms of reduced complexity and cost. Without ISLs, one or multiple satellites are deployed over the same region without peering connections, thus being connected individually to the ground stations. Thereupon, satellites operate in a “bent pipe” mode, where their operation is limited to receiving and retransmitting signals to ground gateways. This approach keeps network intelligence, such as routing, in the ground stations, thus reducing the complexity of computing devices on satellites. Considering the complexity of intra-constellation routing and the long life-cycle of satellites, the simplicity of this design cannot be underestimated.

Constellation Network Typically constellations consist of identical satellites, which may be equipped with a certain number of ISLs (Figure 2.31). This number varies and can exceed the number of adjacent satellites; some constellations introduce multiple ISLs between two satellites in order to enhance connectivity and throughput. However, in general, constellations deploy four ISLs in order to communicate with their neighbors: two ISLs to the former and later satellite within the orbital plane as well as two ISLs to the satellites at the adjacent planes. Thereupon, the underlying form of the constellation network is akin to the bi-directional Manhattan network, with two main differences; first, globally there are two Manhattan networks that are separated by the Orbital seam, and, second, ISLs exhibit variant delays due to the orbiting rotation. The variance can be gradual, caused by the orbital motion of the satellites across the prime meridian, and abrupt, caused by the handovers at the ground stations. Consequently, satellite constellations constitute a constantly shifting networking terrain which challenges the efficiency of fundamental network functions, such as routing.

Megaconstellations

What is meant by the term “Megaconstellations” is the new wave of large constellations of LEO satellites providing ubiquitous broadband connectivity on ground. These initiatives follow the not so successful waves of LEO constellations of the 90’s. These were the Teledesic [95],

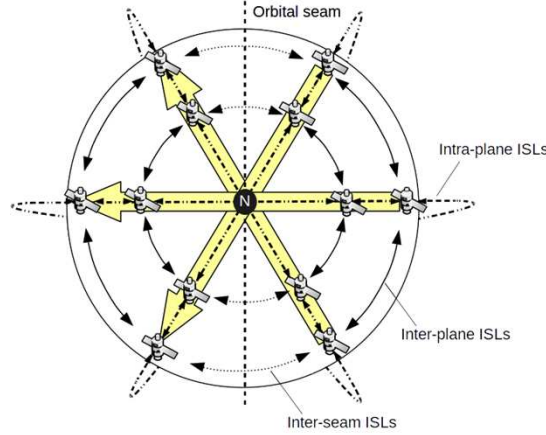


Figure 2.31: Network of star constellation as seen from the North Pole. Arrows depict bidirectional intra- and inter-plane ISLs, as well as inter-seam ISLs.

the Skybridge [93], the Celestri [80], the Iridium [70] and the Globalstar [105], to mention a few. Most of these systems never got off the drawing boards, except perhaps for a couple of demonstration satellites and of course, Iridium and Globalstar which however needed to go into Chapter 11 bankruptcy protection and reorganization before becoming successful. In the recent years, a number of new challengers have come up and proposing new constellations deploying hundreds or even thousands of LEO satellites, i.e the megaconstellations. The business and service models that they follow is often quite different from the ones of the 90's initiatives and as such they all claim a prosperous future. Some examples of megaconstellations are OneWeb [84], LEOSAT [28] and SpaceX [2], although the number of proposals is actually much higher. The main technical characteristics of such constellations are described in Table 2.1 below.

Some satellite renderings are also available for both OneWeb (Figure 2.32) and LEOSAT (Figure 2.33). From these pictures, it is interesting to note the following considerations:

- OneWeb satellite (to be manufactured by Airbus) uses an array of Ku-band antennas to generate a set of 16 user beams whose

	OneWeb	LeoSat	SpaceX
Number of satellites and altitude	648 (882 with spares) 1200 km	78-108 1400 km	4000 1100 km
Satellite mass and power	150 Kg	430 kg, 2.4 kW	100-500 Kg, ?
Payload	Bent pipe, no ISL, 16 user beams in Ku band, feeder link in Ka-band	10 Ka-band steerable user antennas, 2 steerable GW antennas, 4 optical ISLs	Ku band user links ?, Ka band (above 24 GHz) ?
Markets	Broadband, Mobility	Broadband, Mobility	Broadband, Backhaul
Overall claimed throughput and cost	5 Tbps	8 Tbps	1 Tbps
User throughput and latency	50 Mbps, 20-30 ms	1.2 Gbps, 20-30 ms	? Mbps, 50 ms
Beginning of service	2020	2020-2021	2020

Table 2.1: Summary of the characteristics of some megaconstellations systems.

peculiar shape is illustrated in Figure 2.34. According to [84] each beam contains a single forward link carrier of 250 MHz and 6 reverse link carriers of 20 MHz each. The “Venetian Blind” shape of the beams is claimed to be necessary in order to maximize the capacity and availability within the ITU regulatory power flux density limits.

- The rendering of the LEOSAT satellite, to be developed by Thales Alenia Space, clearly shows the 10 Ka band steerable user link antennas and the 2 GW antennas, plus the 4 optical terminals which are used to interconnect the satellite constellation according to the pattern shown in Figure 2.35.

The deployment of the megaconstellations poses a number of challenging technical, business, regulatory and operational issues. In the following, only a subset of the major technical issues is mentioned.

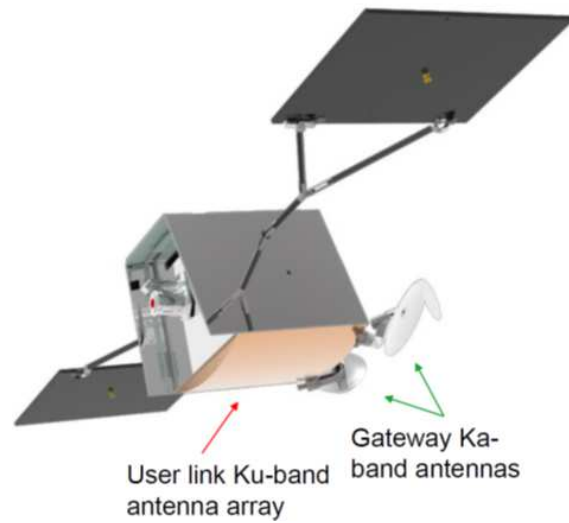


Figure 2.32: OneWeb satellite rendering[84].



Figure 2.33: LEOSAT satellite rendering[12].

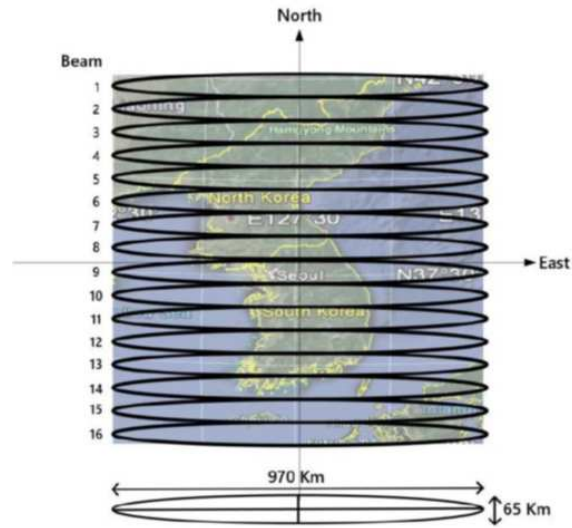


Figure 2.34: OneWeb satellite user link beams[84].

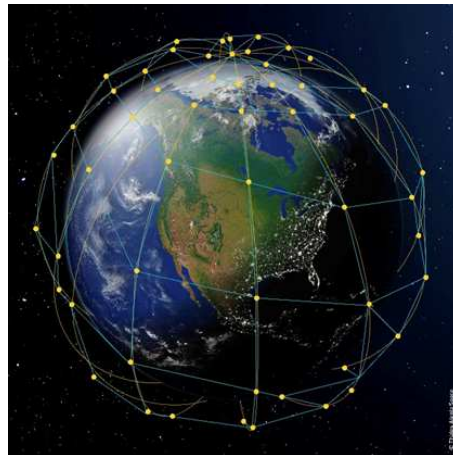


Figure 2.35: LEOSAT constellation showing ISL links[2].

1. Interference with the GEO satellites re-using the same Ku or Ka band. Due to the reuse of the Ku or Ka- band frequency used by GEO telecom satellites, some megaconstellation initiatives will have to deal with interference to and from broadcast and broadband satellites in the geostationary arch. A trivial solution would be to switch-off the satellites around the equator. This solution would result in a relative large unserved geographical area unless some countermeasures are taken like the OneWeb’s “progressive pitching” technique [14].
2. Space debris. The megaconstellation initiatives have the potential to generate a large amount of additional space debris in low earth orbits. The growing population of debris raises fears of more collisions creating yet more debris, thus possibly generating a chain reaction effect [43]. To date, further analysis and discussions are taking place to further understand the issue and possibly reduce its detrimental effect to the space environment.
3. Constellation resource allocation. Due to their intrinsic world-wide coverage, with megaconstellations the issue of optimizing the resources over the areas with actual traffic demand poses additional technical challenges. Indeed, a constellation with no flexibility means could end up illuminating areas with no traffic for a large amount of time thus limiting the effectiveness of the associated business case. Flexibility means may go from a simple selective switch-off of the unused satellites to save battery life, to the exploitation of ISL in order to partially use a subset of satellites as relays, and finally to more complex techniques including active on-board antennas and on-board processors.

2.2.2 High Altitude Platforms (HAPs)

High-Altitude Platforms (HAPs) are aerial platforms that carry communication relays and operate in a quasi-stationary position at altitudes between 15 and 30km¹. They were initially introduced as cost-

¹These days there are also drones or Unmanned Aerial Vehicles (UAVs) to consider; they could play a similar role to HAPs, but at a lower altitude and typically

effective satellite replacements, but lately are regarded more as complementary to satellite networks mainly for providing cheap and agile last mile solutions. Due to the low altitude orbit and the low cost deployment (compared to satellites), several diverse aircrafts can instantiate the HAP concept. For instance, the platform may be an airplane or an airship (or balloon) that can be manned or unmanned and tethered or untethered. The three main requirements posed for HAPs is stability, energy efficiency, and portability. An aircraft must exhibit stability during “flight” for maintaining optical links, energy efficiency for providing long-term services, and easy relocation for adapting to diverse network needs. The most attractive advantages of HAPs originate from their lower flight altitude, which offers sufficient coverage while being quite close to the Earth stations. The proximity of the communication peers suggests cost-efficient operation, as deployment and repair are cheaper and deployment is agile as they are flexible and with fast repositioning. Moreover, HAPs offer typically higher performance communication than satellites because the shorter HAP-terrestrial link introduces less delay and attenuation (and lower error rate). In addition, HAPs have lower power consumption compared to satellites as the communication link can be established with smaller antennas. Finally, the quasi-stationary operation of HAPs does not require handovers and therefore avoids complex networking operations, such as dynamic routing and session continuity during mobility, that are typical in LEO and MEO orbiting satellites. In general, HAPs represent a special point in the “coverage-latency” tradeoff; they fly at significantly lower altitudes than satellites, thus providing higher capacity and lower delay links than satellites with the same budget, but also present relatively limited coverage. Thereupon, HAPs are considered as appropriate solutions for the last mile problem, i.e., delivery directly to the customer’s premises, since it’s an inexpensive broadband solution for areas with increased customer population. For instance, HAPs are proposed for covering cities as second or backup link in order to relieve the load of

with moving profiles. One could imagine swarms of UAVs providing coverage in some area, or following another mobile platform, e.g., a large ship, and including links to a satellite and inter-UAV links.

terrestrial networks or to enhance service availability. Finally, a promising application of HAPs is the delivery of broadband connectivity during an emergency situation or after a disaster, as HAPs provide larger coverage than most terrestrial antennas and allow easier relocation or repositioning than satellites, thus offering fast assistance.

HAP Networking

Due to the coverage limitations of individual HAPs, a reasonable number of aircraft is required for covering large regions. The operation of multiple platforms (e.g., UK requires 5 to 10 platforms for ensuring full coverage), increases the deployment costs and the complexity. Thereupon, the majority of HAP implementations consider one aerial platform that offers service on a regional basis, for instance complementing cellular and satellite networks. There are, however, some studies discussing constellations of HAPs for extending coverage and/or capacity. In the first case, several regionally scattered HAPs with switching payload capabilities exploit inter-platform links to form an aerial network with improved coverage and significant independence to terrestrial infrastructure. Routing of such networks is fairly simpler than for satellite constellations since HAPs are quasi-stationary thus offering a stable topology and avoiding frequent handovers between peering HAPs and, also, between HAPs and ground stations. In the second case, multiple HAPs are deployed in the same area in order to enhance the delivered capacity. Those configurations provide incremental roll-out allowing system expansion by exploiting spatial discrimination techniques.

Integrated HAP-SAT Networks

The integration of HAP and satellite networks is an increasingly popular research topic as it is expected to improve the quality of satellite based services. Given the novelty of this research area, the majority of studies does not delve into the technical integration of satellite and HAP constellations, but, instead, focus on the interoperability of three abstract infrastructural layers: the terrestrial, the HAP, and the satellite. According to this approach the HAP is considered a wire-

less middle step that breaks the satellite-terrestrial downlink into a satellite-HAP and a HAP-terrestrial link (and the same holds for the terrestrial-satellite uplink). Based on the HAP placement and degree of intervention in the terrestrial-satellite link, the HAP-SAT integration may be considered symmetric or asymmetric. In the symmetric design (Figure 2.36), the Earth-SAT link breaks into Earth-HAP and HAP-SAT on both the uplink and the downlink, i.e., the traffic from the satellite as well as from ground stations is first sent to the HAP. This scheme introduces an Earth-HAP link with much lower delay and attenuation, thus addressing specifically the most error-prone part of a satellite link, which is the one closer to the Earth. In the new Earth-HAP link, adaptive techniques, such as adaptive modulation, coding and framing, are more effective than in a traditional satellite link because the propagation of the feedback channel in the former is many times faster than in the later. Similarly, link-layer error-recovery mechanisms provide greater performance gains to e.g. TCP rate control when they are applied at the shorter Earth-HAP link. For instance, the HAP can improve the error-recovery by implementing an Automatic Repeat reQuest (ARQ) protocol that exploits the smaller delay of the Earth HAP link, thus allowing faster error-detection and retransmission. In addition, the HAP can host a Performance Enhancing Proxy (PEP) that takes advantage of the satellite link fragmentation so as to apply two TCP connections that address the individual characteristics of each sub-link, thus enhancing the overall performance².

In the asymmetric design (Figure 2.37), the HAP splits only the return satellite link, i.e., the traffic from the satellite is sent directly to the ground stations, but in the opposite direction, it is first sent to the HAP. This design shares some gains with the symmetric approach such as traffic filtering: the HAP forwards to the satellite only the link-level packets that are received error-free thus avoiding retransmissions through the satellite. Additionally, being the “access point” to the satellite backhaul where numerous users are attached, HAPs can host a cache that can reduce redundant transfers over the satellite link, thus saving satellite capacity. The most fitting scenario of the asym-

²See e.g. [107] for a general discussion of TCP performance over wireless links.

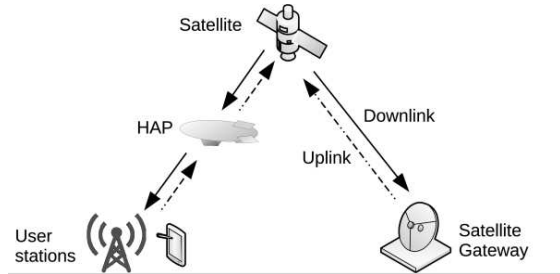


Figure 2.36: Symmetric design of HAP-SAT integration. Optionally a second HAP can be exploited at the Satellite-Gateway link.

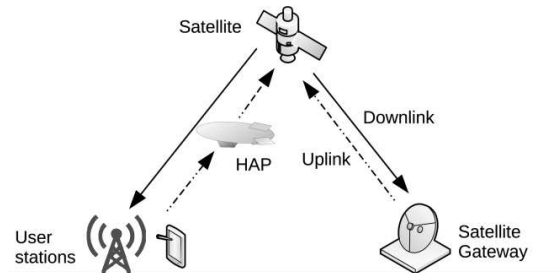


Figure 2.37: Asymmetric design of HAP-SAT integration. Optionally a second HAP can be exploited at the Satellite-Gateway link.

metric scheme assumes numerous mobile users that exploit the faster return channel for the control plane of receiver-driven communication, while the transfer of high data volumes is achieved through a satellite, thus relaxing the capacity requirements of the HAP segment.

2.3 Networking Challenges

Complex, diverse and novel satellite topology and interconnection architectures pose new challenges to networking technology, which has been developed mostly independently and unaware of satellite issues (with some notable exceptions). Even though some satellite systems and notably constellations (e.g., Iridium[70]) have been designed to operate in isolation, independently of terrestrial networks (and actually to compete with them and challenge them), we believe that it's the integration of satellite with terrestrial networks that can provide viable

advantages and opportunities. However, there are various challenges to address, which we touch upon below.

2.3.1 TCP performance

TCP is the dominant transport protocol in the current Internet (TCP/IP) architecture. TCP offers sender-driven reliable delivery and high capacity utilization while also providing network congestion control. An essential part of TCP is its congestion control algorithm that provides fast and agile congestion detection and avoidance[49]. TCP's congestion control introduces an Round-Trip Time (RTT) based approach to estimating available path capacity assuming that RTT increases are caused by congested routers that cannot cope with the overwhelming rate of incoming packets. Thus, TCP increases the rate of transfer while data is timely delivered (probing the network for additional capacity) and decreases it upon the detection of late or lost packets. Finally, RTT estimation relies on Acknowledgment (ACK) packets that are sent back from the receivers in order to report the successful delivery of a data segment.

Even though TCP is the most widely adopted and studied transport protocol in the Internet, its RTT-based rate control is known to perform poorly for 'long' and 'wide' paths, such as satellite links. First, the increase of transfer rate in TCP is inversely proportional to the communication delay, allowing the short paths to grasp bandwidth faster than longer. This issue, also known as TCP-fairness [60], is quite evident in MEO and GEO constellations that exhibit large delays, thus constraining resource utilization. Second, TCP assumes that losses are symptoms of saturated links, hence, in order to prohibit congestion collapse, loss detection triggers reduction of transfer rate. However, satellite networks experience several orders of magnitude higher error rates than wired networks, which can be caused by many operating conditions (mobility, shadowing, weather conditions etc.), different than link saturation, thus leading to unnecessary performance degradation. Numerous solutions have been proposed in order to overcome these problems, such as channel coding to compensate high bit-error rate, special TCP variants that are no longer RTT-based and, also, Performance

Enhancement Proxies[17] that exploit topology awareness to enhance transport layer efficiency. Channel coding introduces methods that address the inherent problems of propagation mediums at the physical or MAC layer. For instance, Forward Error Correction (FEC) introduces redundancy in data transfers that can be used to reconstruct or correct a lost or corrupted packet at the destination. Although channel coding is appealing for most satellite systems for it is transparent to transport layer protocols, its effectiveness is roughly proportional to the added redundancy, therefore expensive resources must be consumed to overcome the higher bit-error rate of satellite links. The second category of solutions consists of satellite-specific TCP variants, such as TCP Westwood[72] and TCP Peach+[8], that unbind rate control from the RTT in order to achieve higher throughput in long satellite links. These designs indeed enhance throughput in satellite links, but fail to deliver those gains in terrestrial networks. Given that the current penetration of satellite networks in the Internet is low, universal substitution of wired TCP variants by satellite ones is not foreseen. An alternative solution suggests that the TCP variant be dynamically selected by the end-users, based on routing intelligence, such as the existence of satellite link in a dissemination path, exposed to the transport protocols. But, IP networks generally conceal path formation from the communication end-points, making such methods impractical. The third category of solutions is based on the introduction of dedicated network entities, called PEPs, which are placed at the satellite gateways and terminals specifically for assisting TCP's performance without requiring changes at the end-users. The most popular PEPs exploit their topological knowledge and techniques such as spoofing and splitting [17], in order to separate connections and rate control into terrestrial and satellite parts. Specifically, spoofing artificially reduces the RTT of the communication by sending earlier "fake" delivery acknowledgements to the TCP senders. Thereafter, the senders perceive lower delays that are solely based on the faster user-PEP network, thus increasing faster their rate of transfer. Similarly, TCP split separates the satellite from the terrestrial network by breaking the initial end-to-end TCP connection into three separate connections: two among the hosts and

the gateways and one among the corresponding gateways, including the satellite path. The third connection uses a satellite TCP variant that meets throughput requirements of the wired TCP variants running at the hosts, thus overcoming the satellite bottleneck. PEPs are widely adopted in satellite networks for providing high performance without requiring changes at the end-users. Nonetheless, PEPs also exhibit significant weaknesses that need to be addressed. For instance spoofing and spitting break the semantics of TCP, thus causing interoperability issues for applications. In addition, spoofing requires symmetric communication, since the delivery of real ACKs detunes TCP's rate control. Furthermore, both splitting and spoofing are not applicable to encrypted transfers, which currently constitute a substantial and increasing portion of global traffic. (PEPs either need to send appropriate ACKs, or create new connections for the same object; hence access to the TCP header is required.) Furthermore, PEPs constitute single-points of failure where considerable amount of work takes place in order to deliver properly data to the destination. All in all, even though PEPs constitute the best solution for satellite TCP performance issues, they still exhibit substantial weaknesses.

2.3.2 Satellite-Terrestrial Network Integration

The deployment of ISLs contributes to the formation of a satellite network that can be independent or complementary to the terrestrial infrastructure. The introduced network can either be completely integrated with the (terrestrial) Internet, or form an Autonomous System (AS), either broadly, or with the narrow technical Internet definition. In the first case the satellite network is transparent to the (IP-based terrestrial) network, thus requiring it to support terrestrial protocols (e.g., RIP and IP-routing). In the second case, the satellite network forms a distinct AS where border gateways running an exterior routing protocol enable communication with other ASs. The main advantage of the integrated scenario compared to the AS approach is that interoperability with the terrestrial network is less complicated; the satellite nodes must be compatible with well-studied and explored terrestrial technologies. However, the integration can result into poor

performance in case the individual characteristics of satellite networks are not considered, for instance acknowledging the long delays of GEO and high bit-error rates of many satellite links. On the other hand, the formation of an AS with non-transparent interior structure can be separating the satellite from the terrestrial network and provide significant gains in terms of performance under some circumstances. The clear separation of the satellite network as well as the acknowledgement of network idiosyncrasies eases the application of network-specific protocols, such as routing protocols that consider delay variance, delivering higher performance. The weakness of the AS approach is that significant overhead must be placed to the resource challenged satellite network in order to consider Internet-scale forwarding tables. The number of forwarding destinations is proportional to the number of Internet domains, thus questioning the design's feasibility within the traditional IP architecture.

Routing

Satellite constellations constitute a continually shifting networking terrain with intense link delay variance due to orbital motion and handovers with ground terminals. In order to enhance the efficiency of routing, novel algorithms that acknowledge the constantly changing link delays are proposed. These designs exploit the predictability and periodicity of network formation, which repeats itself after certain time, and introduce periodic routing phases with differences in the path selection process. For example the Dynamic Virtual Topology Routing method[104] studies the constellation topology in consecutive discrete time intervals $[t_0 = 0, t_1], [t_1, t_2], \dots, [t_{n-1}, t_n = T]$, where T is the period, that are chosen so that:

- Over an interval $[t_i, t_{i+1}]$, the topology can be modeled as a constant graph G_i , i.e. link activation and deactivation take place only at discrete times $t_0, t_1, \dots, t_n$.
- The interval $[t_i, t_{i+1}]$ is small enough to consider the costs of individual ISLs as constant over this time interval. The costs of these links could be computed from a function of inputs such as

distance between the satellites, duration before link deactivation, geographic position, or other factors - assigning higher cost to high-latitude ISLs with a short time remaining before deactivation, for example.

As expected, different routing policies and forwarding rules can be applied on a per-period basis in order to adapt better to the shifting network conditions. For instance, paths can be pre-computed on the ground stations and then uploaded to the satellites at the beginning of each period. Thereupon, arises the necessity for a programmable networking interface at the satellites in order to support frequent updates of forwarding rules.

Mobility

Satellites are constantly traveling across their plane, but the constellation seems stationary to the Earth stations as frequent hand-offs to successive satellites take place. In order to hide the mobility of satellites and interface the constellation with the ground stations, the “Virtual node”[74] concept is introduced. This scheme aims at re-establishing the connection and concealing mobility from the routing protocols.

Information for reaching terrestrial constellation users, such as routing tables and channels used, is state that relates to a particular region of the Earth. This state is maintained in a fixed position relative to the surface of the Earth, at a “virtual node”. A virtual node resides at one satellite at any given time according to the topology of the constellation. In addition, given satellites hosting different virtual nodes, a virtual network of information is created that transfers state according to the constellation movements. This function can be easily supported and simplified through Network Function Virtualization (NFV) discussed later.

2.3.3 Handover in Multi-Gateways HTS Systems

Closely related to the integration of terrestrial and satellite networks is the need for advanced handover techniques in multi-gateways HTS systems. In this scenario, the feeder links are expected to be operated in

EHF frequency band or even in free-space optics, whereby link outages are not infrequent. As a consequence, gateway handover procedures as well as proper frequency re-allocation are necessary to provide service continuity. In this respect, a particular attention has to be drawn on the management of QoS, which can be severely affected because of the additional latency introduced during the handover procedure itself and packet losses resulting once moving from one feederlink to another. In more words, the former reflects the fact that re-routing functions need to be implemented on the “upstream” Internet routers, hence possibly introducing some delays because of the due data forwarding functions. Obviously, this performance impairment may also occur in conjunction with jitter increase, which is particularly undesired in streaming applications.

As to the latter (i.e., packet losses), these can occur because feeder link outage events are forecasted with a non-negligible estimation error and therefore can cause packet losses during the feeder link switch-over unless handover procedures are started with large anticipation, ultimately leading to underutilisation of satellite channel capacity. Moreover, packet losses can be also observed afterwards once the traffic flows managed by the gateway on outage are all handled by the other gateways of the satellite network. These losses are determined by network congestion events, because the feederlink capacity could be not sufficient to accommodate all the traffic flows.

These two points are tightly related to overall QoS management, whose design has been only partly addressed by the scientific community so far. In particular [82] develops a theoretical framework to assess the performance of HTS systems, where network coding strategies are applied to compensate possible packet losses.

Another open research point relates to the simultaneous use of multiple feeder links in order to minimise the occurrence of packet losses during handover events or to deploy more efficient load sharing mechanisms when more than one feeder link can be used to forward data to users (provided that they can receive data from different satellite terminals in different beams). These cases are preliminarily explored in [46], where the use of MPTCP combined to network coding and the

related performance gain is explored. More detail about the potentials offered by MPTCP to satellite networks are further developed in the next chapter of this monograph.

Finally it is worth noticing that proper control of HTS feeder links opens up the door to orchestration concepts applied to the overall satellite network, by taking advantage of Software Defined Networking (SDN) and NFV principles, which are further elaborated in the next subsections. In particular, the advantages coming from the implementation of SDN controllers in the ground segment of satellite network are firstly explored in [16] and further highlighted in [79], where the need for efficient coordination mechanisms is further emphasised to proper transport data throughout HTS satellite systems.

2.3.4 SDN/NFV

SDN constitutes a new architectural paradigm for realizing network functions such as routing, load balancing, etc. SDN is characterized by two basic principles that favor agile network management: first, SDN decouples the control and user plane of the networking equipment and, second, it logically centralizes the network intelligence (i.e. the control plane), thus abstracting the underlying network infrastructure (and improving the chances for applying optimizations). Software defined networking is built upon SDN-enabled switches, which realize the forwarding functionality, and an SDN-controller that undertakes routing decisions. Specifically, the switches receive real-time forwarding rules by the controller via the control plane and apply those rules to the incoming traffic, which constitutes the user plane. Overall, SDN is known to facilitate programmability and simplify management functions via the abstraction of the underlying network infrastructure. NFV is a novel form of function provisioning. The goal of NFV is to decouple network functions from proprietary hardware, making it possible to run such functions in general-purpose commodity servers, switches, and storage units, which could be deployed in a network operator's data center. Unlike SDN, NFV does not necessarily introduce any architectural change into network functions; it embodies an architectural concept that exploits process and node virtualization so as to uncouple functionality

from location. However, NFV usually relies on an SDN-enabled network that realizes connectivity to a transitioning network structure, i.e., it exploits programmable routing and forwarding for adapting to dynamic function provisioning. For instance, SDN can reroute the user plane traffic to the new location of a virtual firewall. All in all, the combination of NFV and SDN simplifies the management of networks, supports agility and resource dynamism, and enhances the utilization of physical resources by orchestrating the placement of multiple virtual functions over a common pool of compute, network, and storage resources. Thereupon, SDN and NFV deliver gains such as lower OPEX, greater flexibility (i.e. easily scaling network resources up and down), and easier network management.

Integrated SDN/NFV-SAT Networks

SDN and NFV are expected to enhance the performance of satellite communications and offer seamless satellite-terrestrial network integration with lower CAPEX and OPEX costs. This is based on the flexibility that SDN/NFV can offer to the “immutable” satellite network infrastructure that exhibits expensive and complex (re)configuration. In the following, we describe the infrastructure of an integrated SDN-SAT network and then discuss several use case scenarios that illustrate the aforementioned gains. Typically SDN/NFV-enabled satellite networks base the usage of SDN switches as gateways in hubs and establish an SDN-controller at each hub. The SDN-controllers are collocated with Network Management Center (NMC) and Network Control Center (NCC) of the satellite network, receiving monitoring information that is used for inter-hub path provisioning and satellite terminal configuration. Finally, an SDN enabled backbone network is frequently assumed for the interconnection of the gateway networks. The overall design concept is shown in Figure 2.38

QoS Routing QoS routing constitutes a popular technique for increasing the performance of networks by setting measurable constraints with regard to the dissemination path, such as delay and capacity [22]. For instance, time critical applications, such as VoIP and teleconference,

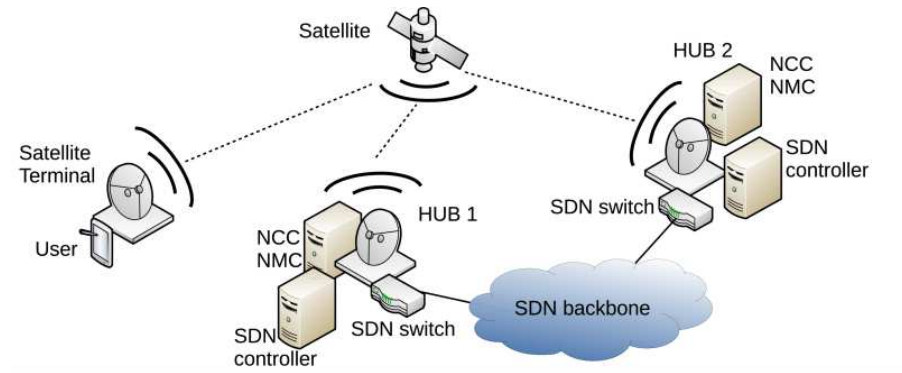


Figure 2.38: SDN/NFV-enabled satellite network.

are known to place demanding requirements about end-to-end delay. This leads to a DiffServ policy, according to which the network operator must not serve these applications over a satellite link when a terrestrial alternative exists. Exploiting SDN, which offers (centralized) controller based decisions routing, the network operator can force serving non-time critical flows over the satellite link, thus relieving the load of the terrestrial infrastructure without violating the agreed QoS constraints. The flow characterization can take place at the SDN-switch located at the gateway, where the Type of Service (ToS)[83] field of the IPv4 header, or the Traffic Class field (TC) of the IPv6 header, could mark the individual requirements of the service or class, or other (e.g., implicit), categorization techniques could be applied.

Service Continuity Satellite networks commonly experience performance degradation due to regional weather disturbances. These issues are addressed by shifting the satellite flows away from the problematic hub to another ground station that is not affected by weather conditions. Expectedly, handing over the active flows to a different hub infers several sub-problems that increase the complexity and cost of the transition, such as PEP state migration, connection re-establishment and re-routing to the new hub. SDN provides efficient handover of Earth stations without requiring service re-establishment. Assuming that sta-

tions are SDN-enabled, then the SDN-controller can update their flow forwarding rules in order to shift traffic towards the new hub. In addition, the SDN-controller can configure the SDN enabled backbone network so as to deliver traffic back to the initial hub and limit the transition in the interior of the satellite network. The final step is to transfer the PEP's state about the active TCP connections to the new hub, which is fairly simple if we assume that PEP is a virtual function that can be transferred atomically to any available hardware device. We further elaborate on PEP migration in the following use case scenario.

Middlebox Virtualization Middleboxes are useful network devices that are deployed in order to apply a specific function or task on traffic, other than simple packet forwarding. By exploiting NFV and SDN, networking middleboxes, such as PEPs, NATs and firewalls, can be provided on demand as a service. Relying solely on software that can be migrated from a standard server to another, middleboxes are easily deployed anywhere, anytime, given that the routing plan of the network is updated in order to push traffic to and from the new location of the function.

PEPs are typical instantiations of networking middleboxes that are quite popular for satellite networks. In order to enhance TCP performance, PEPs usually break the end-to-end semantics of the connection and actively participate in the transmission. Thereupon, essential amount of information about ongoing TCP connections is maintained at the PEPs, determining the state as well as the context of the connection. When a satellite terminal handover takes place, all active TCP connections must be re-established at the new hub's PEP, which requires the (transport of the) state of the initial PEP. Along the lines of the NFV paradigm, PEPs are not implemented as dedicated middleboxes, but rather as software that can be run on any available hardware devices. Consequently, assuming that the PEP function can be assigned appropriately, the "dedicated virtual PEP" can be migrated to the new hub, thus continuing to perform the appropriate TCP optimization.

Enhancing Virtual Network Operator Services Opening up satellite communication devices via programmatic interfaces (with a rich set of instructions that goes beyond SNMP capabilities) exposed to second tier operators, coupled with network virtualization, is an approach to (1) realizing quicker automated service provisioning processes, (2) enriching the service catalogue, and (3) enabling a Satellite Communications as a Service consumption model. By applying device virtualization to Satellite Network Operator (SNO) satellite hubs, a virtual hub can be assigned on a per Satellite Virtual network Operator (SVNO) basis. With the guarantees brought by isolation (which is a key feature of network virtualization and applies to data, control, and management planes as well as to performance and security) an SNO can delegate the full control and management of virtual hubs to their customer SVNOs. Therefore, SVNOs can independently enforce their own policies on their satellite virtual networks. A further step can be achieved by introducing programmability, thus enabling a programmable virtual hub assigned to SVNOs. Programmability may concern the control plane (routing, forwarding, and monitoring, as provided by SDN, allowing SVNOs to devise their own customized traffic control schemes, but also the data plane, allowing SVNOs to devise customized packet processing algorithms (e.g., PEP, encryption, etc.). This paves the way for the diversification and enrichment of the services that SVNOs provide.

3

Networking Solutions for High Throughput Satellite Systems

This chapter explores the application of MPTCP and network coding in those satellite-based scenarios, where the simultaneous availability of multiple path calls for more sophisticated routing and overall networking solutions. Moreover, the occurrence of random loss events may degrade overall system performance, whereby additional recovery functions complementary to channel coding implemented at the physical layer should be envisioned. To this end, the use of network coding is also proposed and more details about its adaptation to the specific scenarios here taken as reference are provided along with some performance results.

3.1 MPTCP

3.1.1 Multiple Air Interfaces and Multi-Homing: Definition, Characteristics, and Examples

Multi-homing is the capability of connecting a host or a computer network to more than one network so that different paths can be simultaneously exploited to reach the same server. Hosts are assigned multiple IP addresses, one for each provider. Multi-homing is today a common approach in heterogeneous wireless networks where a mobile device

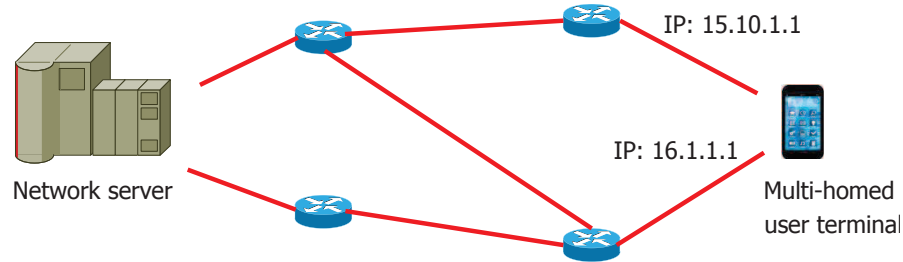


Figure 3.1: Multi-path flows from one source to one destination.

(e.g., a smartphone, a tablet, etc.) has many interfaces and networks to access the Internet, as shown in Figure 3.1. The same device uses distinct IP addresses in different networks (e.g., WiFi and 4G), thus exploiting the inherent diversity supported by a multi-path system [110]. That leads to obvious policy considerations: which interface should be used and when. This could create problems and for sure non-smoothed transitions from the use of one interface to another during a given traffic session (vertical handover). In order to overcome these problems, the idea of multi-path schemes is not to use alternatively these interfaces, but to use them simultaneously to achieve better performance.

We are interested here to study how multipath and multi-homing can be exploited at transport level. In the case of the classical TCP protocol, it can exploit a single path at once: a TCP session will be broken if the device changes the path to reach one server because (for instance) the WiFi connection is not anymore available and there is the need to use a 4G link. When the IP address changes in moving from one coverage area to another, the TCP connection is aborted and restored only when connectivity and a new IP address are available. In this case, the disconnection time could be critical for some applications and quite noticeable to a user.

Novel transport layer protocols have been defined that are able to exploit the multi-homing capabilities of devices to distribute contents towards end-hosts simultaneously using multiple paths. The transport layer session could make use of multiple paths between the two end-

points at any given time so that performance improvements can be achieved: each path could carry data in parallel and congested paths could be avoided, thus preferring faster paths. Sessions could also be more robust. Let us for instance consider a video streaming service using both WiFi and cellular networks; if the user leaves his/her house, the video could seamlessly continue using only the cellular connection. In general, we expect that multi-path protocols can help in performing handovers in mobile IP networks. Other possible application scenarios of multiple paths protocols are data centers (cloud networking), where many servers are interconnected by multiple paths. If the data transport protocol can exploit the simultaneous use of multiple paths, it is possible to dynamically adapt to variable congestion conditions of the paths.

Multi Path TCP (MPTCP) and Stream Control Transmission Protocol (SCTP) are major modifications to TCP that allow multiple paths to be used by a single transport connection [39, 94]. However, SCTP cannot simultaneously use two paths (this would be only possible by adding the Concurrent Multipath Transfer (CMT) extension of SCTP that can really enable load balancing over different paths), while MPTCP makes this possible, thus solving the above problems. MPTCP is based upon the experience gathered in previous works, and goes further to solve fairness issues when competing with other regular TCP flows and deployment issues because of middleboxes in today's Internet. The MPTCP protocol has been standardized by the Internet Engineering Task Force (IETF) in several RFCs, as detailed later. The comparison between the different multipath protocols is provided in Table 3.1. As we can see in this Table 3.1, the main advantage of MPTCP is that it adopts the same APIs as TCP: MPTCP presents itself as classical TCP to higher-layer applications and this is a unique characteristic of this multipath protocol. MPTCP offers the same reliable, in-order byte-stream transport as TCP and is designed to be backward compatible with both applications and network layer. However, MPTCP requires support inside the protocol stack of both endpoints (they need a modified transport layer to support MPTCP and this basically entails changes in the operating systems).

Table 3.1: Comparison of Multipath Protocols.

Basic Features	MPTCP	STCP	CMT-STCP
Compatibility with TCP APIs	Yes	No	No
Compatibility with middle-boxes	Yes	No	No
Simultaneous transmissions via multi-paths	Yes	No	Yes
Organization of data transfer	Byte-oriented	Data chunks	Data chunks

3.1.2 MPTCP Basic Characteristics

MPTCP is an evolution of TCP with three basic goals: (*i*) it has to provide a higher throughput than TCP (*ii*) without taking more resources than necessary and (*iii*) taking as much load as possible from the most congested paths. MPTCP has to use multiple paths maximizing resource usage by splitting the traffic in many sub-flows; the paths could carry data in parallel, and congested paths should be avoided in favor of better paths. MPTCP identifies multiple paths by means of multiple IP addresses at end-hosts.

MPTCP manages the use of sub-flows (activation of sub-flows, disconnection of sub-flows, etc.), the scheduling of traffic among sub-flows (and related paths) and the reordering of the packets at destination before delivering them to the application.

MPTCP starts a sessions activating one end-to-end path (*primary path*). If extra paths are available, MPTCP sub-flows are created on these paths and are combined with the existing session, which continues to appear as a single connection at the application layers of both ends [15]. Each sub-flow that is generated works as a common TCP flow with its congestion window and related state variables. Suitable congestion control schemes at the MPTCP level have been envisaged to control and to balance the growth of the congestion windows (cwnds) on the

paths to account for congestion, differences in Round-Trip Time (RTT), etc. We will discuss later about these schemes. The receiver has only one global receiver window shared between the set of the established sub-flows.

MPTCP uses a flow-level sequence number and a sub-flow level sequence number. There is a mapping between these two sequence numbers that are imbedded in each transmitted packet by means of TCP Options; these options are located at the end of the (MP)TCP header. This mechanism allows segments sent on different sub-flows to be correctly re-ordered at the MPTCP level of the receiver (reordering buffer in the socket). It is considered that Selective ACKnowledgment (SACK) is used at the sub-flow level to achieve better efficiency. A per-sub-flow Retransmission TimeOut (RTO) is adopted to recover from congestion events that TCP fast retransmit/fast recovery cannot handle. The MPTCP protocol stack is depicted in Figure 3.2.

If during the lifetime of an MPTCP session, an interface (and the corresponding path) disappears, the affected host announces this event, so that the peer can remove sub-flows related to this address.

The design of MPTCP has been influenced by many requirements, and, in particular, it as to achieve compatibility with both application layer and network layer, as explained below.

- **Application layer compatibility** implies that applications that today run on top of TCP should continue to work without any change over MPTCP. This implies (as already shown) to use the same APIs as those of TCP.
- **Network layer compatibility** implies that MPTCP must operate over any Internet path where TCP operates (provided that end-hosts are able to support it). Many paths on today's Internet include middleboxes, as detailed below. Unlike IP routers, middleboxes work at transport layer and can use and modify some TCP header fields of connections. The impact of middleboxes is further discussed in the next sub-section.

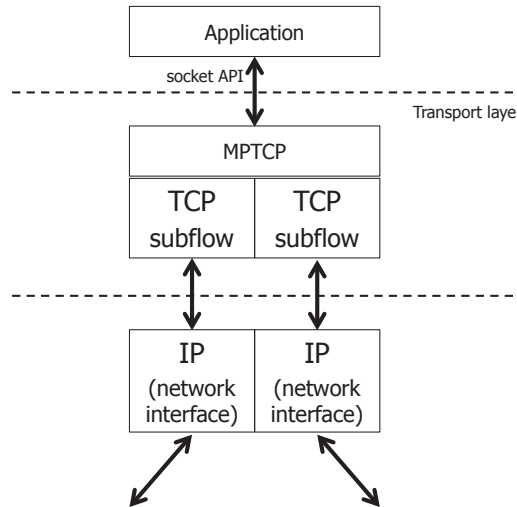


Figure 3.2: MPTCP protocol stack (multi-homed end-host).

3.1.3 Impact of Middleboxes

In the traditional Internet architecture, network devices operate at network or lower layers, with the layers above network layer used only at end-hosts. Today, this layering concept and distinction no longer reflects the reality, since we have a proliferation of middleboxes in the network: middleboxes interpose working at transport layer, sometimes even completely terminating transport connections, thus leaving the application layer as the only real end-to-end layer.

The network compatibility goal requires that the multipath extension to TCP has compatibility with the today's Internet, so that MPTCP sub-flows can traverse predominant middleboxes, such as firewalls, Network Address Translators (NATs), Performance-Enhancing Proxies (PEPs), content engine, streamer, etc. A picture describing the many middlebox types is given in Figure 3.3.

Middleboxes are routers that impose some sort of constraint or transformation (i.e., header modifications) on network traffic passing through them. NAT boxes hide an entire network behind a translation layer that changes the IP address and the port of a connection. Some

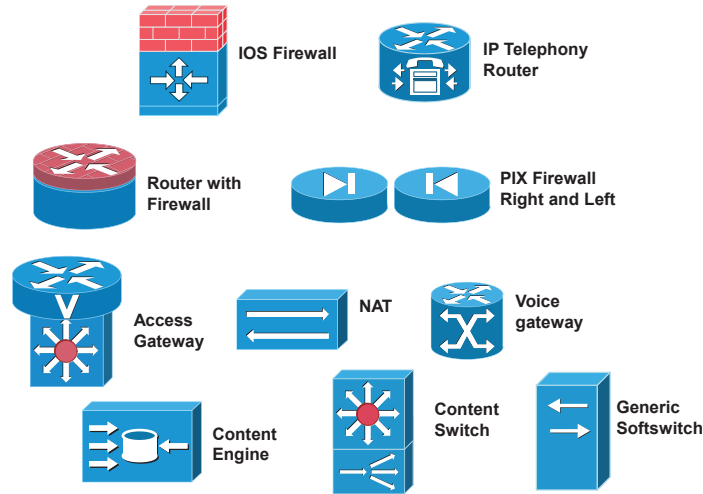


Figure 3.3: Different types of middleboxes.

boxes (like PEPs) will acknowledge data traversing them, well before data arrive at the real destination in order to improve the goodput. Some routers will drop packets with unknown options. Moreover, firewalls will kill connections with holes in the sequence number stream. The list of potential problems is quite long. In addition to the above, if we wanted to adopt a completely new transport layer protocol we could discover that its packets will not cross the Internet. Much of the routing infrastructure assumes that TCP and UDP are the only possible transport-layer protocols; anything else has low probability to work.

The challenge is to design TCP extensions for MPTCP that can safely traverse middleboxes. In order to have that middleboxes be transparent to MPTCP as much as possible, the adopted approach is to transport MPTCP overhead by means of TCP Options [40], a variable-length optional field in the TCP header. MPTCP option has ‘Kind’ equal to 30, reserved by the Internet Assigned Numbers Authority (IANA), ‘Length’ (variable) and the remainder of its content begins with a 4-bit ‘Sub-Type’ field, for which IANA maintains a sub-registry, called “MPTCP Option Subtypes”, under the “TCP Parameters” reg-

istry. Sub-type fields are defined as shown in Table 3.2 according to RFC 6824 [40]. An exemplary TCP option, called Data Sequence Signal (DSS) option, is shown in Figure 3.4 and is used to describe data ACK and data sequence mapping.

3.1.4 Combination of Sub-Flows with MPTCP

MPTCP allows multiple sub-flows to be set up for one MPTCP session. Different operation modes are available for MPTCP as follows:

- **Native MPTCP:** Two MPTCP endpoints establish and use all sub-flows corresponding to the available addresses/port numbers. This mode is used to improve the throughput by means of resource pooling.
- **Active/Backup MPTCP:** Two MPTCP endpoints activate multiple sub-flows, but only a subset of them is actually used in parallel for data transfer. MPTCP endpoints can use the MP_PRIO signal to change the priority of each sub-flow.
- **Single sub-flow MPTCP:** Two MPTCP end-points use just one sub-flow and when a failure is observed, an additional sub-flow is activated so that traffic is forwarded via the newly activated sub-flow.

An MPTCP session starts with an initial sub-flow (primary sub-flow), which is similar to a regular TCP connection. After the first MPTCP sub-flow is set up, additional sub-flows can be activated. Each additional sub-flow is similar to a regular TCP connection, using SYN handshake and FIN tear-down. Rather than being separate TCP connections, all sub-flows are bound into an existing MPTCP session. Data for the connection can then be sent over any of the active sub-flows that has the capacity to take them. The decision on which sub-flow to use for the next packet is a function of the scheduler at the MPTCP layer. The scheduler should decide how much data has to be set via each sub-flow. Different alternatives are available depending on the operating system.

Table 3.2: Sub-Type used by MPTCP Options (prefix 0x is used for numeric constants with hex representation).

Value	Option Symbol	Option Name
0x0	MP_CAPABLE	Multipath Capable
0x1	MP_JOIN	Join Connection
0x2	DSS	Data Sequence Signal (data ACK and data sequence mapping)
0x3	ADD_ADDR	Add Address
0x4	REMOVE_ADDR	Remove Address
0x5	MP_PRIO	Change Sub-flow Priority
0x6	MP_FAIL	Fallback
0x7	MP_FASTCLOSE	Fast Close
0xf	(PRIVATE)	Private Use within Controlled Testbeds

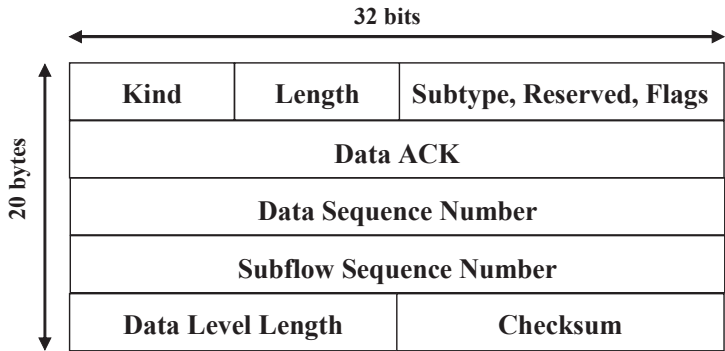


Figure 3.4: MPTCP DSS option, one of the options in Table 3.2 (each option has its own format).

For instance, we can consider a round robin scheduler or a scheduler depending on the RTT value of each path. The MPTCP scheduler has also to deal with a granular allocation: basically, the scheduler decides which sub-flows to service and verifies how much data can be transmitted comparing the congestion window of the sub-flow with the current receiver window value.

The application has not to take care of which path and related sub-flow has the best performance at any instant; MPTCP handles that for the application, trying to privilege the sub-flows with better performance. MPTCP can work when both endpoints are multi-homed (in this case, more sub-flows are opened between all pairs of IP addresses), or even in the case when both endpoints are single-homed (different sub-flows use different port numbers and can be routed differently by multipath routing in the network).

Let us consider the establishment of an MPTCP connection. Let us assume that a multi-homed end-host selects its WiFi interface to open a connection. First, the device sends a SYN segment to the server to start a three-way handshake. The only difference with respect to a classical three-way handshake is that the SYN segment contains the MP-CAPABLE TCP option to show that our device supports MPTCP. This option also contains a cryptographic 64-bit key that is chosen by our device. If the receiving server supports MPTCP, it will add the MP-CAPABLE TCP option and its key to its SYN-ACK reply. This is the only time these keys will be sent in clear; future sub-flows will identify the connection by means of 32-bit “tokens”, obtained by sender and receiver by means of cryptographic hashes of these keys. Our device completes the handshake phase by sending an ACK segment (which must also carry the MP_CAPABLE option) so that the first (primary) path of a multipath session is established: client and server can exchange TCP segments by means of the WiFi access. The exchange of keys in the MP_CAPABLE handshake allows the definitions of tokens that can be used to authenticate the endpoints when new sub-flows are added assuming multiple interfaces are available. A new sub-flow via another path is added by sending a SYN packet to a new interface of the destination with the MP_JOIN option and using

the destination token. This is fundamental to identify which MPTCP session has to be joined.

Naively, our device connected via WiFi could simply send some of the packets over a second interface; however, most ISPs will drop these packets, since they would have the source IP address of another network, the WiFi one. Firewalls and similar middleboxes expect to see a SYN packet before they can admit data packets. Then, the only possible solution is to perform a full SYN handshake on this new path before sending any packets; this is the approach adopted by MPTCP, where there is the need of a three-way handshake phase to activate each path.

An important characteristics of MPTCP, especially useful in the case of smart-phones, is that the set of sub-flows that are associated with an MPTCP connection is not fixed. Sub-flows can be dynamically added and removed from an MPTCP connection throughout its lifetime, without affecting the byte-stream transported. MPTCP also implements mechanisms to add and remove new IP addresses even when an end-point is behind a NAT. If the smart-phone connects to another WiFi network, it will receive a new IP address. At that time, it will open a new sub-flow using its newly allocated IP address and tell the server that its old address is not valid anymore. The server will now send data only towards the new address. These options can help the smartphone moving through different wireless connections without stopping its MPTCP connection.

Each segment sent by MPTCP has two sequence numbers: the Sub-flow Sequence Number (SSN) inside the regular TCP header, and an additional (global-level) Data Sequence Number (DSN) inside a TCP option. This solution ensures that the segments sent on any given sub-flow have consecutive sequence numbers and have no conflicts with middleboxes. MPTCP can then send some data sequence numbers on one path and the remainder on the other path; old middleboxes will ignore the DSN option that will be used by the MPTCP layer at the receiver to reorder the byte-stream before it is delivered to the receiving application. The data stream as a whole can be reassembled (in-order delivery to application layer) through the use of the data sequence

mapping components of the DSS option, which defines the mapping from SSN to DSN.

MPTCP provides a connection-level acknowledgment, acting as a cumulative ACK (data successfully received in order) for the connection as a whole. This is the “Data ACK” field of the DSS option that is based on DSN (see Figure 3.4). Instead, the sub-flow-level ACK is based on SSN and acts analogously to TCP SACK, considering that there may be holes in the data stream at the sub-flow level.

3.1.5 MPTCP Congestion Control Schemes

MPTCP is an extension to the classical TCP, which allows users to spread their traffic across potentially disjoint paths. MPTCP discovers the paths available to a user, establishes the paths, and distributes traffic across these paths by means of separate sub-flows. Multipath-capable flows should be designed so that they shift their traffic loads from congested paths towards uncongested paths, so as to achieve a better use of Internet capacity. In effect, multipath congestion control means that the sender carries out a task that is normally associated with routing, namely moving traffic onto paths to avoid congestion. The three design criteria for MPTCP congestion control can be expressed as follows:

- **Design goal 1:** To be fair with respect to regular single-path TCP flows,
- **Design goal 2:** MPTCP should use paths efficiently (optimizing the aggregate throughput),
- **Design goal 3:** To perform (at least) as well as single-path TCP.

MPTCP allows *resource pooling*, meaning that it allows to aggregate resources on different paths to improve the goodput. Among the solutions that effectively pool resources, we should prefer those that balance traffic between the sub-flows depending on the degree of congestion of the paths; this is what we call *to equipoise the traffic load* among the paths.

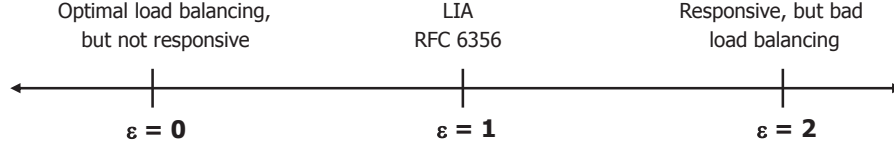


Figure 3.5: Load balancing (i.e., to equipoise the MPTCP goodput among the paths) versus responsiveness (i.e., capability to quickly adapt to changes in the available capacity).

MPTCP congestion control algorithm has to achieve a trade-off between optimal load balancing of the traffic flows among the paths (avoiding oscillations of traffic also known as flappiness) and responsiveness (i.e., the capability to react promptly when the congestion level changes) [88]. Assuming to have two paths only with different characteristics (e.g., different RTTs), the two extreme cases are: (i) we send the traffic only on the best path; (ii) the congestion control schemes on the two paths behave like single and independent TCP flows (in this case, since the packets of the two flows have to be reordered at the MPTCP layer of the receiver, the slower flow delays also the packets of the other flow: the worst path of the two is determining the performance). To study all the possible congestion control cases, including the two extreme ones above, a parameter $\varepsilon = [0, 2]$ is used, where $\varepsilon = 0$ corresponds to the first case (also called **Fully Coupled algorithm** for optimal load balancing [54]) and where $\varepsilon = 2$ corresponds to the second case (also called **Uncoupled TCP flows**, because there is no attempt to balance the traffic loads among the paths). The interest is to find an intermediate optimal ε value combining the congestion control of the two paths and improving the total goodput. Figure 3.5 shows the consequences of using different ε values: in general, a lower ε value yields a better traffic load balancing among paths, moving the traffic away from congestion; instead, higher ε values allow to achieve more responsiveness.

Let w_r and RTT_r be the cwnd and the estimated RTT on path $r \in R_u$, respectively, where R_u is the set of all paths available to user u . We denote by l_r a smoothed estimation of the number of bytes sent

over path r between the last two losses seen on r . Note that $1/l_r$ can be considered as an estimate of the loss probability on path r , p_r [57].

The idea is that the laws to increase the cwnds w_r on the different paths be related each other to avoid that one MPTCP flow be unfair against other (classical) TCP flows. The general ‘semicoupled’ congestion control scheme based on parameter ε can be stated as follows, referring to TCP NewReno congestion avoidance phase [106]:

- For each ACK on sub-flow r , cwnd w_r is updated as:

$$w_r = w_r + \frac{a^{2-\varepsilon} w_r^{1-\varepsilon}}{w_{total}^{2-\varepsilon}}. \quad (3.1)$$

- For each loss on sub-flow r , decrease w_r as follows:

$$w_r = w_r - \frac{w_r}{b}, \quad (3.2)$$

where a and b are parameters that can be set to mimic the behavior of classical single-path TCP (i.e., $a = 1$ and $b = 2$) and where

$$w_{total} = \sum_{r \in R_u} w_r. \quad (3.3)$$

The idea behind this congestion control algorithm is to transmit over path r at a rate proportional to $p_r^{-\frac{1}{\varepsilon}}$, where p_r is the packet loss probability over this path [106].

The reaction to a packet loss does not depend on ε and then this is common to all cases. Basically $b = 2$ is used as in TCP NewReno. Let us consider below the reaction to the ACK for the different cases. We are interested to select ε in the range $[0, 2]$ to see if it is possible to achieve both resource pooling and traffic balance. The throughput of an MPTCP flow reduces when ε decreases from 2 to 0.

The case $\varepsilon = 0$ corresponds to the **Fully Coupled algorithm**: the traffic is sent only over the best paths. The congestion control scheme with the Fully Coupled scheme is as follows:

- For each ACK on sub-flow r , cwnd w_r is updated as:

$$w_r = w_r + \frac{a^2 w_r}{w_{total}^2}. \quad (3.4)$$

The coupled congestion controller, inspired by [54], achieves optimal load balancing in static networks with all paths having similar RTTs. However, in practice this algorithm is not responsive (it fails to detect free capacity in dynamic settings) and ‘flappy’ (when there are multiple good paths with similar loss rates, it randomly flip traffic among them). Coupling moves traffic away from the more congested path until loss rates are equalized on the paths. But losses are never exactly the same, so traffic oscillates randomly between the two paths.

The case $\varepsilon = 2$ corresponds to using **Uncoupled TCP flows** on each path. With the uncoupled algorithm, the congestion window of each sub-flow behaves like for a single standard TCP connection. The congestion control scheme in the uncoupled case is described below:

- For each ACK on sub-flow r , cwnd w_r is updated as:

$$w_r = w_r + \frac{1}{w_r}. \quad (3.5)$$

Oscillations are reduced with this algorithm, but it does not balance congestion among sub-flows: this algorithm is not very good at moving traffic away from a congested path.

The intermediate case $\varepsilon = 1$ corresponds to the **Linked-Increases Algorithm (LIA)** that avoids flappines (as it occurs in the fully coupled case) and achieves a compromise between optimal resource pooling and responsiveness. LIA works as follows:

- For each ACK on sub-flow r , cwnd w_r is updated as:

$$w_r = w_r + \frac{a}{w_{total}}. \quad (3.6)$$

LIA allocates to path r a window w_r , which is proportional to the inverse of the loss probability $1/p_r$, so that the total MPTCP throughput is equal to the rate that a regular TCP user would get on the best path, that is $\max_{r \in R_u} \left\{ \sqrt{\frac{2}{p_r}} \times \frac{1}{RTT_r} \right\}$ [106]. Parameter a controls how aggressively to increase cwnds, hence it controls the overall throughput. The optimized choice for a is based on fairness requirements (i.e., improve throughput, do no harm, and balance congestion). According to [88], LIA a parameter is designed to achieve that the total MPTCP throughput is equal to the throughput of a single TCP flow running on the best path. Then, parameter a results as follows [54]:

$$a = w_{total} \frac{\max_{r \in R_u} \left\{ \frac{w_r}{RTT_r^2} \right\}}{\left(\sum_{r \in R_u} \frac{w_r}{RTT_r} \right)^2}. \quad (3.7)$$

This formula for a obviously requires that the congestion control scheme uses RTT measurements for the different paths by means of timestamps.

A variant of LIA has been defined where the cwnd increase is controlled by taking a minimum. This new congestion control scheme is called “RTT compensator”:

$$w_r = w_r + \min \left\{ \frac{a}{w_{total}}, \frac{1}{w_r} \right\}. \quad (3.8)$$

The minimum has been included in the cwnd increase so that the increase cannot be larger than that of the uncoupled case (stand-alone TCP flow). The minimum is useful when the paths have unbalanced RTTs. Then, combining (3.7) and (3.8), we achieve the following law to increase the cwnd at each ACK with LIA - RTT compensator algorithm:

$$w_r = w_r + \min \left\{ \frac{\max_{r \in R_u} \left\{ \frac{w_r}{RTT_r^2} \right\}}{\left(\sum_{r \in R_u} \frac{w_r}{RTT_r} \right)^2}, \frac{1}{w_r} \right\}. \quad (3.9)$$

We have seen that MPTCP is not responsive (it fails to detect free capacity in dynamic conditions) and is highly flappy (MPTCP would use one path almost exclusively for a while, then flip to another path,

then repeat). The solution to these issues is to use the LIA algorithm defined in RFC 6356. This algorithm improves throughput (the total MPTCP throughput with LIA is better than that of TCP over the best path). Moreover, LIA is not more aggressive than a single TCP connection. However, LIA has a poor performance in balancing congestion among paths [55, 56]. Moreover, upgrading some TCP users to MPTCP can reduce the throughput of classical TCP users without any benefit for the upgraded users. MPTCP can also be excessively aggressive towards TCP users when single-path and multipath users share resources. This implies that LIA is not Pareto-optimal. In order to provide both responsiveness and load balancing we need the new congestion control scheme, called **Opportunistic - LIA (OLIA)**. LIA achieves a trade-off between responsiveness and load balancing. Instead, OLIA simultaneously provides responsiveness and load balancing [57].

OLIA adapts cwnd increases as a function of (i) the number of transmitted bytes since last loss to make it responsive and non-flappy and (ii) the RTTs of paths to compensate for different RTTs. OLIA is implemented in the Linux kernel 3.0.0. If both paths have similar characteristics, OLIA uses both of them (non-flappiness); instead, if the second path is congested, OLIA uses only the first path. Hence, OLIA satisfies the design goals of LIA specified in RFC 6356 (i.e., to improve throughput, do no harm other TCP flows, and to balance congestion among paths), thus providing optimal congestion balancing and fairness in the network: OLIA is Pareto-optimal.

The OLIA congestion window management for each path r is as described below referring only to the action corresponding to each ACK (since the reaction to a packet loss is always the same, that is to halve the congestion window):

$$w_r = w_r + \frac{\frac{w_r}{RTT_r^2}}{\left(\sum_{r \in R_u} \frac{w_r}{RTT_r}\right)^2} + \frac{\alpha_r}{w_r}, \quad (3.10)$$

where α_r coefficients are used to consider those cases where the best paths (i.e., set of paths r that have the maximum value of l_r^2/RTT_r) are different from those with the largest congestion windows; the details on how α_r coefficients are obtained can be found in [55, 56]. Note

that the first term in (3.10) depending on RTT is meant to optimize the congestion balance among paths compensating for different RTTs, while the second term is used to guarantee responsiveness and non-flappiness of OLIA.

3.1.6 MPTCP Applied to Mobile Satellite Systems

MPTCP has many areas of applications, including large scale intercontinental backbone and data centers with meshed connectivity (cloud networks). Here, we consider in Figure 3.6 some possible satellite scenarios where MPTCP can play a significant role. In all these scenarios, we have mobile users that exploit multiple paths to connect via satellite thus achieving a suitable degree of diversity.

Mobile users with multiple air interfaces of the same satellite system (cases b, c, e in Figure 3.6) : In this case, we exploit the overlap in the coverage of two beams of two satellites or different paths via the same satellite so that two air interfaces can be used simultaneously on the mobile terminal. The capability to exploit multiple paths is useful to improve capacity (we can expect that using simultaneously the two paths we can achieve better capacity than using a single path) or to better support the mobility when the mobile terminal is connected via satellite and moves away from a beam towards another (in the transition area, the terminal can benefit from using simultaneously both paths, thus improving the capacity during the handover phase). This could be the case of the flight scenario b, when a plane connected via satellite is reaching the border of the current beam and can use two paths via the old beam and the new one to improve goodput and reliability.

Vertical handover using a satellite path and a terrestrial one in an integrated system (case d in 3.6) : This is a situation similar to the previous ones, but here the multiple paths use two different wireless technologies having different characteristics in terms of RTT, information bit-rate, packet loss rate. MPTCP can allow a smoother behavior

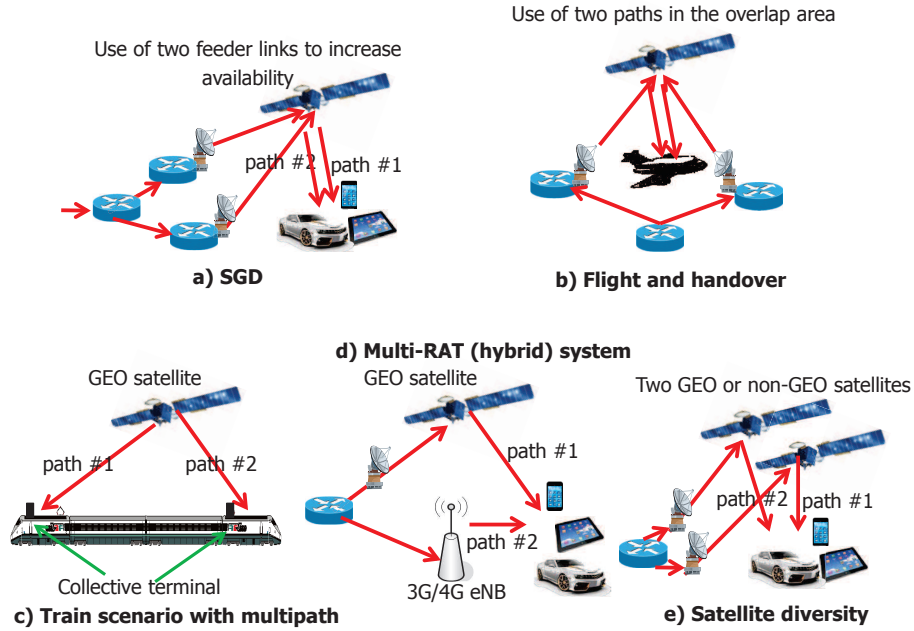


Figure 3.6: Multi-path scenarios suitable for MPTCP use in mobile satellite systems.

of the vertical handover process between terrestrial and satellite systems.

SGD (case a in Figure 3.6) : In this case, multiple paths are possible via different GWs to reach the satellite and then end users; the GW feeder link adopts Q/V/W bands to increase capacity, but at these frequencies the link is very sensitive to atmospheric effects so that there is the need to adopt GW handover and rerouting schemes when a GW experiences bad atmospheric conditions. MPTCP can activate the new path via a new GW in clear sky conditions before the old GW is disconnected so that the remaining packets that cannot be delivered through the old path can be seamlessly transmitted via the new path.

PEP-based architectures (especially in case c of Figure 3.6 and in all cases where a collective terminal can be exploited to support a PEP at the user side) : This is a special case, since we refer to the adoption of two PEPs to isolate between them the portion of the network where MPTCP is used, so that end-hosts are transparent to MPTCP (they do not need an operating system upgrade to support MPTCP). A PEP can be installed at both an intermediate router and at a collective terminal to support the split approach.

3.1.7 Effects of Unbalanced Paths

We consider here the performance implications of MPTCP in the case of unbalanced paths in terms of propagation delays, capacity, and packet losses. As already explained, MPTCP with LIA may suffer from cases with unbalanced paths. Since the packets have to be delivered in order by MPTCP layer, the slowest path causes a delay also for the other paths that have to buffer data waiting for the missing packets from the slowest path. This entails some issues: the socket buffer of MPTCP has to be of adequate size in order to avoid packet losses. More critical conditions are caused by links affected by packet losses because we have to wait for them to be recovered via retransmissions. MPTCP can suffer if the path with the lowest RTT is also the one with the highest Packet Loss Rate (PLR). This occurs because LIA privileges the path with the lowest RTT for traffic delivery; hence, if this path is also the one with the highest PLR, then the MPTCP goodput of these two paths can also be lower than the goodput of the best path of the two.

3.1.8 MPTCP RFCs and Implementations

MPTCP is defined by some IETF RFCs, among which we can consider the following ones:

- RFC 6181 - Threat Analysis for TCP Extensions for Multipath Operation with Multiple Addresses
- RFC 6182 - Architectural Guidelines for Multipath TCP Development

- RFC 6356 - Coupled Congestion Control for Multipath Transport Protocols
- RFC 6824 - TCP Extensions for Multipath Operation with Multiple Addresses.

MPTCP implementations are already available on a number of operating systems. MPTCP is available for Linux, FreeBSD, Android, iOS7 and OS X Yosemite. The currently-available MPTCP implementations are detailed in reference [4]. An MPTCP kernel implementation is available at the link in [3].

3.2 Network Coding and its Applications to Satellite Networking

In contrast to source coding, where only data sources can send encoded packets, network coding allows intermediate nodes in a network to re-encode packets (i.e., several received data packets can be encoded into a single code word) before sending them instead of adopting the simple store-and-forward routing approach [7]. Network coding remained a theoretical concept, until the COPE system was presented in [53], a new architecture for wireless mesh networks where routers mix (encode) packets from different sources to increase the information content of each transmission. Then, in 2007, network coding was used by the Microsoft Secure Content Distribution (MSCD), a Peer-to-Peer (P2P) file distribution network. Network coding operation is realized at the packet payload level and is performed at the source as well as possibly at some intermediate nodes. This coding strategy can allow several advantages, such as more efficient and more reliable communications and an increase of throughput.

It is known that network coding improves multicast capacity, reduces bandwidth requirements for data replications in distributed cache systems (big data applications), and reduces file-downloading latency in P2P file sharing systems.

Let us consider a system that acts as an information relay, such as a router, a node in an ad-hoc network, or a node in a P2P distribu-

tion network. Traditionally, when a node is forwarding an information packet destined to other nodes, we have simply to repeat the packet, applying the Kirshhof's law "What goes in, goes out" (Figure 3.7a). Instead, network coding allows that the output of a node be a function of a mixture of the input information: the node combines a number of packets it has received into one or several outgoing (encoded) packets, as shown in Figure 3.7b. Moreover, if compared to traditional single-path routing protocols, opportunistic routing in wireless mesh networks can improve unicast throughput by utilizing the shared wireless broadcast medium. In particular, all neighboring nodes of a transmitter may overhear the data packet and may cooperate in forwarding the packet to its destination. The use of network coding in this scenario allows both combining the packets reaching the destination from several neighboring nodes without taking care of the exact packet sequence and being robust to packet losses. NC operates at the level of binary representation of the packets. In particular, encoding can be performed on single bits [XOR operations over Galois Field (GF) of size 2] or on groups of m bits, that is *symbols* (over GF of size 2^m); when $m = 8$, symbols are bytes and this is a quite convenient case for implementations.

The main steps in the evolution of coding up to network coding are depicted in Figure 3.8.

Network coding was originally proposed to achieve multicast data delivery at the maximum transfer rate in multicast networks with single source. Network coding can be used when there is the need to reliably transmit a number of messages to a set of users such as in a reliable multicast system [45]. In traditional Automatic Repeat reQuest (ARQ), each single lost message has to be retransmitted. By means of network coding techniques multiple users can recover different lost packets using just one retransmission [44]. Since these first examples, network coding has found applications in many areas, such as computer networking, distributed storage, P2P networks, security, etc. A survey on possible application examples of network coding in computer networks is given in [73]. In terrestrial wireless mesh and peer-to-peer networks, the benefits of network coding are very well understood in terms of throughput and robustness [53]. Moreover, there are many network coding tech-

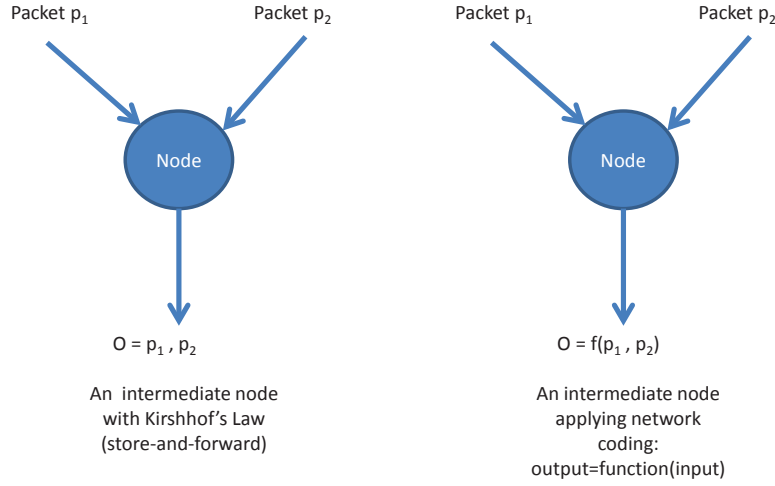


Figure 3.7: Intermediate nodes without network coding [on the left, case (a)] and with network coding [on the right, case (b)].

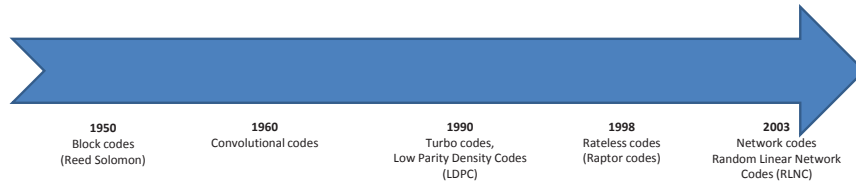


Figure 3.8: Evolution of coding from the beginning to network coding.

niques that can be adopted in satellite systems, as shown in [102]. The work in [101] studies load balancing in a multi-beam satellite system with network coding. In [47], network coding is used to save bandwidth in a bidirectional communication between two satellite terminals.

Network coding application areas are:

- Reliability,
- Congestion control,
- Multi-path routing,
- Security.

3.2.1 Packet-level FEC Versus Network Coding

Physical layer Forward Error Correction (FEC) methods are used to correct the errors at the bit/symbol level. However, in the case that there are too many errors, FEC methods might not be able to correct the errors and the packets will be dropped since the MAC-level Cyclic Redundancy Check (CRC) control fails; this process would cause packet erasures. In this case, packet-level **FEC codes** (e.g., fountain codes) as well as network coding could help to recover packet losses. MAC-level packet correction is different from PHY-level bit correction, because it deals with straight packet losses, not with unpredictable bit errors. Packet-level FEC can be applied to packets on top of MAC layer on a link-basis.

The objective of erasure correction is to recover from the loss of entire packets. The transmission of K input packets is complemented by that of $N - K$ extra (redundancy) packets. If at least K packets out of N are correctly received, then all K input packets can be retrieved. The redundancy overhead is $(N - K)/N$. There are scenarios where the reduced error rate is very valuable. When multicasting data towards large groups, even a small individual error rate per recipient may result in large retransmission rates for the whole group so that the use of redundancy will result in dramatic efficiency gains. In the case of long transmission delays (as in the satellite case), the use of redundancy helps maintaining the delivery delays within acceptable limits even in the presence of errors, because in most of the cases the decoding process should be able to directly recover the losses without the need of retransmissions. A brief survey about the most common FEC approaches is provided below.

Fountain codes are a class of erasure FEC codes with the property that a potentially limitless sequence of encoded packets can be generated from a given set of source packets (rateless erasure codes) and the original source packets can ideally be recovered from any subset of the encoded packets of size equal to or only slightly larger than the number of source packets. The terms ‘fountain’ and ‘rateless’ refer to the fact that these codes do not exhibit a fixed code rate. Luby

Transform (LT) codes and Raptor codes are examples of Fountain codes [67]. In particular, Raptor and RaptorQ codes are detailed in RFCs 5053 and 6330 [68],[69], respectively.

Raptor codes use simple XOR operations over $GF(2)$, while **RaptorQ codes** use a mixture of $GF(2)$ and $GF(2^8)$ operations. The Raptor/RaptorQ encoder is applied independently to each source block. Each source block is identified by a unique integer Source Block Number (SBN). Each source block is divided into a number, K , of source packets, each with size T bytes. Each source packet is identified by a unique integer Encoding Symbol Identifier (ESI) [103]. According to RFC 5053, the header overhead of Raptor codes is: a 16-bit SBN and a 16-bit ESI. According to RFC 6330, the header overhead of RaptorQ codes encompasses 8-bit SBN and 24-bit ESI. Raptor (RaptorQ) codes can be used for encoding together up to $K = 8,192$ (56,403) packets in a source block, having a total of up to $N = 65,536$ (16,777,216) packets per encoded block. IETF RFC 5053 adopted Raptor codes for broadcast/multicast file delivery applications. Systematic Raptor codes have been adopted into multiple standards, such as the 3GPP Multimedia Broadcast/Multicast Service (MBMS) standard for broadcast file delivery and streaming services, the DVB-H IPDC standard for delivering IP services over DVB networks, and DVB-IPTV for commercial TV services over IP networks.

Reed-Solomon codes are another example of packet-level FEC according to the definition in RFC 5510. Encoding and decoding procedures employ arithmetic in $GF(2^m)$. Each packet is subdivided into words of length m bits. The packet stream that is actually transmitted consists of FEC groups containing both data packets and check packets used for reconstructing lost data packets. Data packets have fixed size and check packets have the same length as data packets. A Reed-Solomon code is specified as $RS(n, k)$ with m -bit symbols. In real applications m is equal to 8, 16, or 32. This means that the encoder takes k data symbols, each of m bits, and adds parity symbols to obtain an n -symbol codeword. There are $n - k$ parity symbols.

Encoder and decoder perform operations in $\text{GF}(2^m)$. Given a symbol size m , the maximum codeword length n for a Reed-Solomon code is $n = 2^m - 1$. A Reed-Solomon decoder can correct up to t erroneous symbols or up to $2t$ erasures in a codeword, where $2t = n - k$. An erasure occurs when the position of an erroneous symbol is known. RS(255,223) is a typical example of Reed-Solomon code.

Important packet-level FEC applications deal with reliable multicast transmissions, where packet data have to be sent (without errors) to a group of receivers. The work in [38] presented a reliable multicast framework where repair requests and retransmitted packets are always multicast to the whole group. Multicast packets may be affected by errors at different stages of transmission; the more receivers we add to the group, the larger is the probability that some of them will lost any given packet. In order to avoid to transmit the whole set of packets again, the idea is to adopt a packet-level FEC, where encoded packets are sent to a multicast group in blocks of N . If a receiver is able to obtain K correct packets out of N it is able to decode a block [48]. Otherwise, it has to signals back how many new coded packets it needs to complete the decoding. A random access protocol is used to send the feedback. The sender in receiving the first feedback can retransmit a number of extra redundancy packets that is exactly the number of packet losses indicated in the feedback message. These extra redundancy packets are useful not only for the considered terminal, but also for all other terminals with losses because of the packet-level FEC. This technique has been adopted by the NACK-Oriented Reliable Multicast (NORM) protocol [6].

There are many characteristics that differentiate network coding from a common packet-level FEC. In particular, only network coding can allow the **recoding of an encoded packet** at intermediate nodes. The recoded packet can be used together with first-coded packets to decode a block. Moreover, other specific characteristics of network coding are: encoding applied at network or at transport layer, inter-flow network coding. Figure 3.9 shows the operation of a generic network

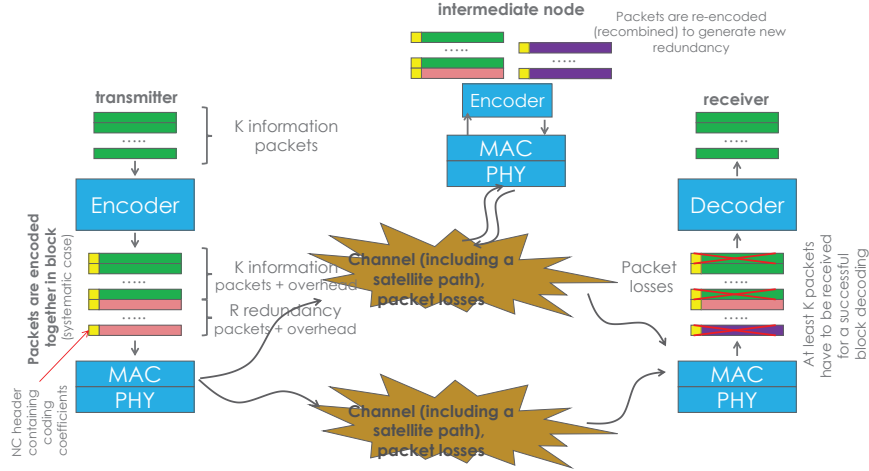


Figure 3.9: Packet-layer encoder/decoder at sender and receiver sides (in this example, network coding is a shim layer between IP and MAC).

coding scheme for unicast flows, where the encoder works on a block of K source packets that can be successfully decoded at the receiver only if enough redundancy packets are generated to mask errors experienced on the different links (at MAC layer, a packet is considered lost if and only if the layer 2 fails to correct it according to the CRC). In other words, the destination should receive at least K coded packets that are linearly independent in order to successfully decode the block.

The Internet Research Task Force (IRTF) has approved the Network Coding Research Group (NWCRCG) [1] to investigate several areas concerning the application of network coding to the networks, such as: architectural considerations (e.g., control plane, routing plane, transport layer), end-to-end versus hop-by-hop network coding, inter-flow network coding (at layer 3 or below) versus intra-flow, application-layer network coding, security issues and robustness to attacks, packet formats, proactive protection with network coding versus reactive mechanisms based on Automatic Repeat reQuest (ARQ) techniques (trade-off between redundancy and retransmission delays).

A still open issue is, how could network coding be integrated in current networking protocols. In particular, the OSI layer where to

apply network coding is still a subject of research. There are many options on the OSI layer where to apply network coding, that is:

- Between IP and MAC layers (shim layer),
- Between transport and IP layers (shim layer),
- Between application and lower layers.

These network coding options entail different pros and cons and compatibility issues with other protocols, such as IPsec, header compression, etc. For instance, IPsec is applied end-to-end so that basically recoding at intermediate nodes would be impossible if IPsec is adopted. The OSI layer where network coding is performed also implies the position where network coding header (encoding coefficients) is placed in the sequence of the progressive encapsulation process from application down to physical layer. The headers of the different layers are sequenced like this: PHY | MAC | IP | Transport | APP. Note that, with network coding, we have an additional header overhead to be placed properly in the sequence:

- PHY | MAC | **NC** | IP | Transport | APP
- PHY | MAC | IP | **NC** | Transport | APP
- PHY | MAC | IP | Transport | **NC** | APP.

We have **intra-flow (or inter-flow)** network coding when coding is applied over payloads belonging to the same flow (or to multiple flows). This differentiation also depends on the layer where network coding is applied. We can roughly consider that intra-flow network coding is possible when network coding is applied at high OSI layers (i.e., transport level and above); instead, inter-flow network coding can be adopted when network coding is applied at lower layers (network level and below). The decision between these two alternatives is a tradeoff between performance gains and complexity of operation.

Mixing packets from different flows (inter-flow network coding) can be convenient for some applications (like routing schemes), but in other cases it could even lead to a poorer decoding performance (exhaustively mixing data of several flows may decrease the chance a receiver can recover its information).

Note that when NC layer is on top of MAC layer, NC can be integrated with MAC layer encapsulation that is used to transport IP packets in the MAC frame via satellite. Multi-Protocol Encapsulation (MPE) and its more efficient evolution, called Generic Stream Encapsulation (GSE), are two examples of MAC-layer encapsulation. If the air interface is DVB-S2, GSE [or Return Link Encapsulation (RLE)] encapsulation is used for the forward (or return) link.

NC can be performed block-by-block or on a sliding window basis [37]:

- In **block coding**, the original payload sequence is divided in blocks of given size and network coding is performed only applying a transformation over the payloads of a block. A coding block is a set of usually consecutive packets defined by the sender side.
- In **sliding window coding**, given a stream of un-coded payloads, coding blocks are selected on the basis of a sliding window. Coding blocks are partially overlapping. This is a form of convolutional encoding at the packet level instead of at the bit level.

If NC is of the block coding type, the input flow has to be segmented into a sequence of blocks, and encoding and decoding must be performed independently block-by-block. This approach has a major impact on coding and decoding delays, since encoding requires that all the source symbols be available at the encoder. The block creation time also represents the minimum decoding latency. A good value for the block size is necessarily a trade-off between the decoding latency (we cannot decode a block until all the packets of a block are received) and the desired robustness to bridge long erasure bursts, which depend on the block size. On the other hand, a convolutional code associated

to a sliding window removes this decoding delay, since repair symbols can be generated and sent on-the-fly from the source symbols in the current encoding window [29]. The sliding encoding window mode is especially convenient for real-time applications since the decoding time is lower.

In block coding, there is the need to partition the data into multiple blocks and encode only packets in the same block. Let us assume that each packet consists of B bits. When the packets to be combined do not have the same size, the shorter ones are padded with 0 bits. We can interpret q consecutive bits of a packet as a symbol over $\text{GF}(2^q)$ so that each packet consists of a vector of $L = \lceil B/q \rceil$ symbols.

3.2.2 Min-Cut Theorem

In this section, we refer to broadcast and multicast applications where many end hosts need to receive the same set of packets. This is a quite interesting case for satellite systems.

Let us model a communication network as a directed graph $G = (V, E, C)$, composed of a vertex set V and a directed edge set E [25], where the edges in the graph have integer capacity C (the capacity of an edge is the maximum amount of flow that can pass through that edge). One or more edges with unit capacity are considered between the nodes connected by higher rate links. In graph theory, a cut is a partition of the vertices of a graph into two disjoint subsets. An s - t cut is a cut that has source and sink in different subsets; the cut-set only consists of edges going from the source side to the sink side. The capacity of an s - t cut is the sum of the capacity of each edge in the cut-set. The value of a flow is defined as the amount of flow passing from source to destination.

Let us consider a directed acyclic graph $G = (V, E)$ with unit capacity edges, h unit rate sources located on the same vertex of the graph and N receivers. Let us assume that the value of the min-cut to each receiver is h . Then, there exists a multicast transmission scheme over a large-enough finite field $\text{GF}(2^q)$, in which intermediate network nodes linearly combine their incoming information symbols over $\text{GF}(2^q)$, that delivers the information from the sources simultaneously to each re-

ceiver at a rate equal to h . Hence, the maximum flow in a network is dictated by its bottleneck (= min-cut).

In the unicast case, a session (s, t) consists of a single sender s and a single receiver t . Let $r(s, t)$ denote an achievable rate for the communication from s to t in a given (V, E, C) network. An upper bound on $r(s, t)$ is the value of any s - t cut through the network. The minimum of the values of all such s - t cuts, herein designated as $MinCut(s, t)$, is also an upper bound on $r(s, t)$. That is, $r(s, t) \leq MinCut(s, t)$. For undirected graphs with unit-capacity edges, it has been proved that there always exists a set of $h = MinCut(s, t)$ edge-disjoint paths between s and t . Thus, by routing information over this set of h unit-capacity edge-disjoint paths, we can achieve reliable communication between s and t at the maximum possible rate, $r(s, t) = MinCut(s, t)$.

A simple network coding scheme could use XOR (exclusive-or, denoted as “ $\oplus$ ” or simply “+”) over GF(2), i.e., a bit-wise addition between packet data for a quick encoding and decoding scheme. When a node receives packets from two different sources, simply XOR them together, and send them to the next nodes. Decoding is still carried out by means of XOR operations. An example of the above considerations is provided below, referring to the butterfly network in Figure 3.10, where a multicast source (at the top of the picture) has packets A and B to be sent to the two destination nodes (at the bottom). Each link can carry only one packet per time slot. If only classical routing is allowed, then the central link would be only able to carry A or B, but not both by means of a coding operation. Suppose we send A through the center; then the left destination would receive A twice and not B. Sending B poses the same problem for the right destination. Then, classical routing is inefficient because no routing scheme can transmit both A and B simultaneously to both destinations. However, using a simple coding approach, as shown in Figure 3.10, A and B can be transmitted to both destinations simultaneously by sending the XOR of packets A and B (i.e., $A + B$) through the center (see also Figure 3.7). The left destination receives A and $A + B$, and can calculate B by doing the XOR of A and $A + B$. Similarly, the right destination will receive B and $A + B$, and will also be able to recover A by doing the XOR of

B and $A + B$. Then, the coding approach allows to achieve the best throughput for the butterfly network, that is 2 packets delivered to the destinations in one slot time. With classical routing, we would need extra transmissions to deliver both packets A and B to both destinations. The butterfly network has permitted to show how XOR network coding outperforms a classical routing-based approach in terms of throughput.

Let us provide further considerations on the above example. If we cut the network graph into two parts, one containing the source and the other containing a receiver and then record the net flow across the boundary from source to receiver, then the maximum value of the s - t flow is equal to the minimum capacity over all s - t cuts (*max-flow min-cut capacity*). In our butterfly network example, a receiver is able to receive A and B packets in one slot time from the two links for a throughput of 2 packets per slot time (= max-flow min-cut capacity). Therefore, it has been shown that XOR network coding is sufficient to achieve this optimum capacity.

3.2.3 Arithmetic in a Galois Field

In order to understand better how network coding can be carried out using GFs, we need to recall about arithmetic in a Galois field. It is different from integer arithmetic. With a Galois field $\text{GF}(2^q)$, addition, subtraction, multiplication and division operations are defined over the numbers $0, 1, \dots, 2^q - 1$ in such a way that: if a and b are elements belonging to $\text{GF}(2^q)$, then also $(a + b)$, $(a - b)$, $(a \times b)$, and (a/b) are elements of $\text{GF}(2^q)$. An element of $\text{GF}(2^q)$ can be represented by q bits as $(a_{q-1}, a_{q-2}, \dots, a_1, a_0)$ and can also be represented in a polynomial form of degree $q - 1$ over $\text{GF}(2)$ as $A(x) = a_{q-1}x^{q-1} + a_{q-2}x^{q-2} + \dots + a_1x + a_0$. The sum of two elements of $\text{GF}(2^q)$ is obtained doing the XOR sum (equivalent to use the arithmetic sum modulo 2) of bits of same position. In general, operations in $\text{GF}(2^q)$ are performed modulo a certain *primitive polynomial* $P(x)$ over $\text{GF}(2)$ that has degree q with binary coefficients 0 and 1. Note that there are 2^n polynomials over $\text{GF}(2)$ with degree n . A polynomial $P(x)$ over $\text{GF}(2)$ of degree n is said to be irreducible over $\text{GF}(2)$ if $P(x)$ is not devisable by any polynomial over $\text{GF}(2)$ of degree less than n but greater than zero. For instance,

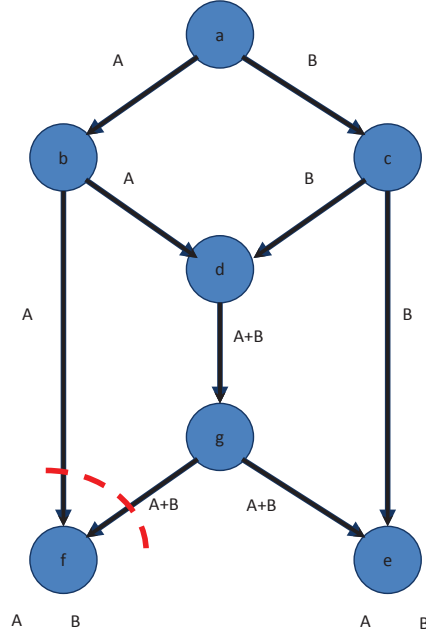


Figure 3.10: XOR network coding applied to a butterfly network for multicast transmissions.

$x^2 + x + 1$ is an irreducible polynomial over $\text{GF}(2)$ of degree $n = 2$. Moreover, $x^3 + x + 1$ is an irreducible polynomial over $\text{GF}(2)$ of degree $n = 3$. An irreducible polynomial $P(x)$ of degree n is said to be primitive if the smallest positive integer j for which $P(x)$ divides $x^j + 1$ is $j = 2^n - 1$. For a given n , there may be more than one primitive polynomial of degree n . There are tables that provide primitive polynomials. For instance, $P(x) = x^3 + x + 1$ is a primitive polynomial over $\text{GF}(2)$ with degree $n = 3$. A good implementation of network coding libraries requires an efficient execution of operations in $\text{GF}(2^q)$.

In the following description on network coding schemes, we refer to block coding techniques.

3.2.4 Linear Network Coding

In Linear Network Coding (LNC), outgoing packets are linear combinations of the original packets, where addition and multiplication are performed over $\text{GF}(2^q)$ using suitable multiplying coefficients. The reason for choosing a linear framework is that the algorithms for coding and decoding are well understood. LNC across the network induces a linear transformation on the source vector message. The network transform characteristic can be then modeled by a system matrix induced by the linear code. The network operates a transformation, that projects the source codebook into the receiver codebook by the system matrix $\mathbf{M}$. The receiver decodes the source messages by inverting such transformation $\mathbf{M}$. The coefficients must be strategically chosen so that matrix $\mathbf{M}$ is invertible.

Assume that a number of original packets $P_1, \dots, P_K$ is generated by one or several sources. In LNC, each packet is associated with a sequence of coefficients $\alpha_{i,1}, \dots, \alpha_{i,K}$ belonging to $\text{GF}(2^q)$ and corresponds to $X = \sum_{j=1}^K \alpha_{i,j} P_j$; coefficient products and summations have to be carried out at the symbol level within a packet [63]. An encoded packet has to contain both the coefficients $\alpha = \{\alpha_{i,1}, \dots, \alpha_{i,K}\}$ in the header and the encoded data $X = \sum_{j=1}^K \alpha_{i,j} P_j$ in the encoded packet payload. The encoding vector α is used by recipients to decode the data by reconstructing the matrix $\mathbf{M}$ of the encoding coefficients.

A special case is when we adopt $\text{GF}(2) = \{0, 1\}$: a symbol is a bit, and the linear combination sent by a node for instance after having received packets P_1 and P_2 is $P_1 + P_2$, obtained as the bitwise XOR operation.

3.2.5 Random Linear Network Coding

In the case of Random Linear Network Coding (RLNC), coefficients $\alpha_{i,1}, \dots, \alpha_{i,K}$ are selected randomly and independently with uniform distribution from the elements of $\text{GF}(2^q)$. Let us assume that a source has to send K packets to a set of receivers. Each native packet is composed of L symbols in $\text{GF}(2^q)$ so that each symbol is composed of q bits. Let $P_i = (p_{i,1}, \dots, p_{i,L})$, $i = 1, \dots, K$, denote the i -th native

packet to be sent by the source and $p_{i,j} \in \text{GF}(2^q)$ be the j -th symbol of the i -th packet. The K source packets arrive at an intermediate node that will generate coded packets from the K original ones. Each coded packet F_i , $i = 1, \dots, K$ can be expressed as:

$$F_i = \sum_{j=1}^K \alpha_{i,j} P_j, \quad \text{for } i = 1, \dots, K, \quad (3.11)$$

where sum and products operations are according to $\text{GF}(2^q)$, as explained later.

Equation (3.11) can be organized in a matrix form as:

$$\begin{bmatrix} F_1 \\ F_2 \\ \vdots \\ F_K \end{bmatrix} = \begin{bmatrix} \alpha_{1,1} & \alpha_{1,2} & \dots & \alpha_{1,K} \\ \alpha_{2,1} & \alpha_{2,2} & \dots & \alpha_{2,K} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{K,1} & \alpha_{K,2} & \dots & \alpha_{K,K} \end{bmatrix} \cdot \begin{bmatrix} P_1 \\ P_2 \\ \vdots \\ P_K \end{bmatrix}. \quad (3.12)$$

If a packet P_i is shorter than L symbols, a series of zeros is appended to it to meet the length L .

If the K coded packets, F_i , are broadcast by the intermediate node, every receiver can try to decode these packets by inverting (3.12) using the Gaussian Elimination (GE) technique to retrieve the source packets. If some transmitted linear combinations (encoded packets) are lost because of the channel, the full block of K source packets cannot be recovered and, for this reason, an additional number of $\delta = N - K$ encoded packets needs to be generated and sent by the intermediate node to compensate for possible losses, where N has to be determined

adequately.

$$\begin{bmatrix} F_1 \\ F_2 \\ \vdots \\ F_K \\ F_{K+1} \\ \vdots \\ F_N \end{bmatrix} = \begin{bmatrix} \alpha_{1,1} & \alpha_{1,2} & \dots & \alpha_{1,K} \\ \alpha_{2,1} & \alpha_{2,2} & \dots & \alpha_{2,K} \\ \vdots & \vdots & \ddots & \vdots \\ \alpha_{K,1} & \alpha_{K,2} & \dots & \alpha_{K,K} \\ \alpha_{K+1,1} & \alpha_{K+1,2} & \dots & \alpha_{K+1,K} \\ \alpha_{N,1} & \alpha_{N,2} & \dots & \alpha_{N,K} \end{bmatrix} \cdot \begin{bmatrix} P_1 \\ P_2 \\ \vdots \\ P_K \end{bmatrix}$$

$$\Leftrightarrow \mathbf{F} = \mathbf{M} \cdot \mathbf{P}. \tag{3.13}$$

A received packet is called *innovative* if it increases the rank of the matrix. If a packet is non-innovative, it is reduced to a row of 0s by GE and is ignored. In order to correctly decode the packets, we need that the matrix $\mathbf{M}$ of the encoding coefficients has rank equal to K so that we can recover the original K packets by inverting matrix $\mathbf{M}$.

Due to the random selection of the encoding coefficients, it may happen that some rows of the matrix are linearly dependent, so that some encoded packets could not contribute to the decoding process to invert $\mathbf{M}$. To reduce the risk of linearly dependent rows, the best approach is to increase the GF size so that encoding coefficients can be selected in a larger set. Therefore, a good choice for the GF size, which will work efficiently is from $\text{GF}(2^8)$ to $\text{GF}(2^{16})$.

The coding/decoding technique of RLNC presented herein is referred to, in literature, as batch coding, which requires that the entire batch (set of linear combinations) is received before being able to decode: the decoder can start decoding to retrieve the block of original packets after having received the packets of the batch.

An encoded packet needs a suitable header to convey the encoding vector: $\{\alpha_{i,1}, \alpha_{i,2}, \dots, \alpha_{i,K}\}$. This requires $K \times q$ bits plus additional 16 bits to contain the progressive ID of the encoded block (i.e., generation ID), the Source Block Number (SBN).

There are three important open issues in RLNC:

- High decoding computational complexity because of the GE method;
- High transmission overhead because of the need to include the coefficients vectors in the header of every encoded packet;
- Linear dependency among coefficients vectors that can reduce the number of innovative encoded packets.

Pseudo Random Network Coding A variant of RLNC is represented by Pseudo Random Network Coding (PRNC). PRNC uses a pseudo random code generated from a codebook row, and includes the index (called ‘seed’) of the codebook row instead of the full coefficients vectors in the header of the encoded packet. Because the coefficients are taken from a codebook, the decoder can recover the coding coefficients of an encoded packet by means of the index in the packet header. PRNC helps to reduce the overhead of packet header with respect to RLNC. The index is put into the coded packet header and does not require any further protection. The drawback is a slightly reduced independence among combinations because the use of the index reduces the number of possible combinations. Moreover, PRNC does not allow recoding at intermediate nodes, while RLNC does it.

Systematic RLNC The traditional RLNC approach is to choose the coding coefficients at random from the used GF, a strategy that is called Dense Network Coding (DNC), since all transmitted (encoded) packets are dense combinations of the original packets. Consider that for a generation of K original packets $K + \delta$ coded packets are sent. If more than δ packets are lost, the matrix of received coding coefficients will not be full rank, and no packets can be decoded. In this case we lost the full block of K packets even if without network coding some of these packets could be delivered at destination; hence, in this case network coding is not so efficient. To avoid this problem, systematic network coding can be adopted [86]. For each generation of K original packets, $K + \delta$ packets are sent, where the first K packets are the original information packets. The packets received of these K can be

delivered to the higher layer, irrespectively of the reception of other packets of the same block, since they are already decoded. The last δ packets (redundant packets) are coded using coefficients chosen at random among the GF elements, except zero since a densely coded packet maximizes the probability that the encoded packet is innovative so that it can be used to compensate for the loss of any original packet.

RLNC Decoding and Complexity The receiver maintains a decoding buffer for storing the coded packets received within a block. The decoding schemes requires to invert matrix $\mathbf{M}$ of the coding coefficients. Matrix $\mathbf{M}$ has size $N \times K$, with $N \geq K$. In order to decode the original K packets, the rank of matrix $\mathbf{M}$ has to be equal to K , the maximum possible. Row elimination methods are adopted to remove dependencies among rows. The Gauss-Jordan elimination method can be used to invert matrix $\mathbf{M}$.

Even working on Galois fields, matrix inversion methods are formally the same as in the classical linear algebra; the only difference is that operations are now on GF. Matrix inversion can be carried out on square matrices. Row echelon form refers to a modified triangular matrix where among other things, $N - K$ dependent rows are the last ones with all zeros. The GE algorithm is used to obtain the row echelon form. Gauss-Jordan elimination is a variant of the classical GE method to transform a matrix to the Reduced Row-Echelon Form (RREF), in which redundant equations are removed and each row contains 0s except for diagonal elements that are 1s. The benefits of RREF is that, with the matrix reduced to an identity matrix, the result vector on the right of the equation constitutes the decoding solution.

RLNC with Gaussian-Jordan elimination has decoding complexity of $O(K^3)$. Larger block sizes could entail long decoding times that could be still acceptable for some data applications, but not compatible with real-time services. Large block sizes could be for instance useful to recover packet losses that span large intervals.

Just to make a comparison with the complexity of packet-level FEC codes, we can consider Raptor and RaptorQ codes that have the advantage that encoding and decoding complexities are $O(K)$.

However, sparse end-to-end erasure correcting codes, such as Luby Transform (LT) and Raptor codes lack RLNC's capability that packets can re-encoded at intermediate nodes so that re-encoded packets and originally-coded packets can be used together for decoding the same block. Even if sparse end-to-end erasure codes have a low complexity, they entails additional delays, because sparsity reduces the impact of a coded packet (i.e., its degree of innovation). Dense coded packets are innovative with high probability, instead sparse coded packets are innovative with low probability [34].

With RLNC, we need to wait for all the encoding vectors to be received in order to decode a block of K packets; this decoding time could be critical for some real-time applications; it would be better that packets be delivered one by one to the application. Thus, instead of waiting for all the packets of an encoded block to arrive, a partial decoding process can be carried out at the reception of each packet. The received packets containing the coefficient vectors and coded data blocks are organized as an augmented matrix in which a transfer unit constitutes a row such that the progressive decoding can be carried out on this matrix. Decoding starts as soon as the first packet of a block is received and re-attempted at each new packet arrival by adding more rows in the matrix. In order to perform this 'progressive decoding' scheme [91], we need to use the Gauss-Jordan elimination algorithm, since it can be performed progressively as encoded packets are being received. This scheme is uniquely possible with RLNC using dense code matrices. Progressive decoding reduces the overall decoding latency when the arrivals of data blocks to be decoded span a long period of time. In this process, the decoding time overlaps with the time required to receive the packets, thus hiding part of encoding and decoding times. However, the decoding process can be completed only at the arrival of K independent encoded packets of a block.

Decoding Failure Probability Let us refer to non-systematic RLNC. An **encoded packet is innovative** for a user if the corresponding encoding vector is linearly independent of all the encoding vectors already received by that user. If we encode K packets into $N \geq K$ packets, we

need at least K independent encoded packets to be able to correctly decode this block without the need of any further action. The condition for the correct decoding of a block is $\text{rank}\{\mathbf{M}\} = K$, so that matrix $\mathbf{M}$ is full rank. The probability that matrix $\mathbf{M}$, $N \times K$, is full rank is equal to $P_{\text{full}} = \prod_{j=0}^{K-1} (1 - Q^{j-N})$, where $Q = 2^q$ denotes the GF size, $\text{GF}(2^q)$. With sufficiently large GF size (for instance $Q = 256$, $q = 8$ bits per symbol), the probability not to have enough independent packets in a block is negligible. More critical is the situation when the communication channel introduces packet erasures. Redundancy is intended to compensate for the loss of packets: let us assume to have ℓ lost packets in a block; then, the received matrix is $\mathbf{M}'$ that is $(N - \ell) \times K$. Considering that $(N - \ell) \geq K$, we can decode the block only if $\text{rank}\{\mathbf{M}'\} = K$ and this occurs with probability $P_{\text{full}}|\ell = \prod_{j=0}^{K-1} (1 - Q^{j-N+\ell})$. Let i denote the number of packets correctly received out of N of a certain block: $N - i = \ell$. Then, in the RLNC case, the probability P_{suc} to successfully decode a block of N encoded packets in the presence of a random erasure channel with packet loss probability p can be expressed as:

$$P_{\text{suc}} = P_{\text{suc}}(N, K, p, Q) = \sum_{i=K}^N \binom{N}{i} (1-p)^i p^{N-i} \prod_{j=0}^{K-1} (1 - Q^{j-i}), \quad (3.14)$$

where we use the binomial distribution to characterize the probability of i packets correctly received out of a block of N packets (memoryless erasure channel).

The decoding failure probability is the complementary of P_{suc} : $P_{\text{fail}} = 1 - P_{\text{suc}}$.

If the decoding of an RLNC block fails, there is the possibility to use a feedback scheme to ask for further redundancy to be transmitted by the sender according to a sort of ARQ scheme (rateless approach); this is possible because RLNC codes are rateless and because the sender holds in a buffer the packets until it receives the confirmation of correct delivery. RLNC without feedback is described in [78]; instead, RLNC with feedback is dealt with in [30].

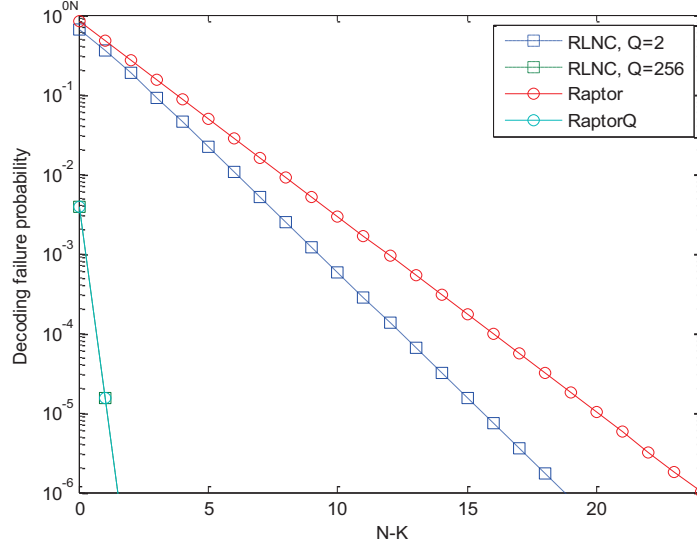


Figure 3.11: Decoding failure probability of different schemes with $K = 5$ and no packet losses.

The decoding failure probability behaviors of RLNC and Raptor codes are compared in Figure 3.11 as a function of the redundancy degree $N - K$, taking also the dependency among linear combinations into account for different Q values. We can see that the decoding performance of both RLNC with $Q = 2$ and Raptor codes is not so good. A better performance can be achieved using RLNC with $Q = 256$ and RaptorQ codes. Performance improves with Q , since it is less probable to find linearly dependent combinations as Q increases.

3.2.6 Instantly Decodable Codes

Real-time applications have strict deadlines and can tolerate some packet losses. Despite this fault tolerance, these applications can suffer significantly from packet losses that cause high jitter, thus for instance impacting the quality of video services. Therefore, it is important to be able to recover packet losses with very low delay and within a very small coding window. This is possible by adopting Instantly Decod-

able Network Codes (IDNC) for loss recovery that allow instantaneous decoding (zero delay).

Even if RLNC achieves better throughput performance, IDNC schemes are interesting because of the following properties. IDNC allows instant packet decodability upon successful packet reception at the receivers. This decodability property allows a progressive recovery of the encoded contents. The encoding process of IDNC is carried out by means of simple XOR operations as compared with more complex operations over large GF as performed in RLNC. This XOR encoding scheme requires a lower packet overhead as compared with the encoding vector overhead required by RLNC. Finally, the decoding process of IDNC is performed using XOR operations, which are less complex as compared with the matrix inversions needed by RLNC.

IDNC techniques shift the computational load to the sender, and allow receivers to perform simple binary XOR-based packet decoding, assuming that the sender usually has larger computational and energy resources than receivers. The sender has to find a strategy to send packet XOR combinations that can convey decodable information at the receivers. A single coded transmission may allow more receivers to decode missing packets by means of simple XOR decoding operations. IDNC approaches are suitable for real-time broadcast applications, where RLNC codes cannot be used because of their long decoding times. For instance, let us assume that the sequence of packets X_1 and X_2 has to be provided to destinations U_1 and U_2 via lossy links. In a first attempt, destination U_1 correctly receives only packet X_1 , while destination U_2 correctly receives only packet X_2 . Without network coding, the source has to retransmit both packets X_1 and X_2 . With IDNC network coding, we can transmit just one encoded packet $X_1 + X_2$, so that both U_1 and U_2 can soon recover both packets.

In the COPE system [53], an opportunistic inter-session network coding scheme is adopted for wireless networks. Each node uses knowledge of what its neighbors have overheard to perform opportunistic network coding so that each encoded packet can be decoded immediately at the next hop. This scheme combines packets in a First-Input First-Output (FIFO) way (as stored in the transmission queue) and

greedily maximizes the number of receivers that can decode the packets in the next time slot.

3.2.7 Applications of Network Coding to Satellite Networks

We provide below a short summary of some network coding applications to satellite networks.

The paper in [65] presents design, test, and implementation of a **P2P on-demand video streaming system with network coding** of the RLNC type. The advantage of network coding is that it allows a perfect coordination, since an arbitrary number of serving peers (referred here to as seeds) can be used to serve the same segment to a receiving peer. In particular, a PRNC approach is used where the index of the codebook row instead of the full encoding vector is included in the header of an encoded packet. In the context of P2P on-demand media streaming, each media segment is divided into blocks. The server sends coded blocks of a segment to its connected peers. When a peer has decoded a segment it can recode the segment and send new encoded blocks to its downstream peers. With Gauss-Jordan elimination, the receiving peer can progressively recover original blocks as it receives coded ones in a ‘pipelining’ fashion, so that the entire segment is immediately playable after the last block is received. With network coding, coded blocks from any peers are equally useful: a peer may download coded blocks of a segment from multiple peers at the same time, without requiring any protocol to coordinate the efforts of these peers. Without network coding, a peer will have to explicitly request particular blocks from a particular server or from serving peers. The purpose of using network coding is to exploit peers upload bandwidth in order to reduce bandwidth needs for the server. As in any streaming mesh, peers have to exchange content availability information each other in order to select appropriate serving peers.

A **TCP / Network Coding (TCP/NC)** source transmits random linear combinations of all packets in a sliding coding window that is related to the congestion window [96]. The sink acknowledges every

degree of freedom (i.e., an innovative linear combination) even if it does not immediately reveal an original packet. The purpose of network coding here is to protect TCP from packet losses. Network coding is integrated at transport layer and this requires a modification of the operating system kernel as well as of the TCP header. Transport layer payloads are stored and encoded according to linear combinations in a sliding coding window. Each coded packet is considered as a degree of freedom, so that packet ordering is not important at the receiver. When enough packets are received to decode the linear combinations, decoded data are delivered to the transport layer. Whenever the source is allowed to transmit, it sends a random linear combination of all packets in the coding window. For each received packet, the receiver computes the rank of the decoding matrix; the receiver acknowledges every degree of freedom received and not every packet (i.e., the receiver acknowledges every independent encoded packet received). This entails to compute the rank of the decoding matrix for every packet received and this could be heavy from the computational load standpoint for the receiver; the complexity to be managed is cubic with the number of packets currently in the decoding matrix. The encoding scheme used in this work is similar to the on-the-fly encoding scheme recently proposed in [29].

The availability of large spectrum portions in Q, V, and W (EHF) bands has paved the way to the design of future satellite networks, able to meet users' demands in terms of high-quality and high data-rate services for emerging applications such as 3D, 4K and Ultra High Definition TV. As already explained, the exploitation of EHF frequency bands for the feeder link (between earth station and satellite), introduces formidable technical challenges due to severe propagation impairments, thus requiring a careful design of the overall system. To address this problem, techniques have been proposed to improve the service availability on the basis of the SGD concept. The SGD technique takes advantage of the spatial diversity existing between earth stations (also called GWs), using a GW handover scheme so that the traffic flows transported by the feeder link going

to suffer from outage are rerouted towards a more favorable GW in terms of meteorological conditions (clear sky) and traffic congestion. The main idea developed in [81] is to use **network coding for a GW handover scheme** to allow some flows to use the GW feeder link until disconnection so that we can improve the efficiency. Before the disconnection occurs, when the feeder link SNR goes below a certain ‘attention’ threshold, we start to reroute the traffic towards another GW in clear sky conditions and we send the redundancy packets of the current buffer contents (inter-flow network coding applied just below IP layer) towards the other GW to protect the traffic flows on the feeder link that is going to experience disconnection. With this technique it has been possible to prove that higher efficiency can be achieved for SGD systems.

Let us refer to an SGD scenario where multiple paths (multi-homing) are available for mobile satellite users exploiting MPTCP. In order to counteract the effects of ON/OFF channels experienced by mobile users, we consider to adopt network coding to protect each MPTCP sub-flow [46]. In order to exploit diversity the encoded block of each sub-flow is sent in part via one path and in part via another path. The decoder decodes the blocks coming from the same encoder and delivers the decoded packets to the MPTCP layer. **Network coding is applied at a shim layer below IP layer** (inter-flow network coding). It has been shown that this technique can improve the TCP goodput in comparison with similar multipath schemes in the literature.

A final example deals with **network coding applied to reliable multicast traffic** via satellite with the cooperation of a terrestrial segment. In this case, we have the transmission of encoded data packets (RLNC) from a source via satellite to multiple receivers. Due to the satellite channels, some packets can be lost so that the redundancy of RLNC can be used to recover the losses. A terrestrial complementary component (WiFi/3G/4G/5G) can be used to cooperate for the transmissions of recoded packets. Because of the property of network

coding, there is not the need to retransmit the specific missing packets for each receivers, but just extra coded packets of the same block can be used by all the receivers to recover the block. A single extra redundancy packet can be useful to recover the losses of different packets for different receivers, so that network coding represents a very efficient approach for reliable multicast applications.

4

Future Trends in Satellite Networking

This chapter discusses Satellite Communication (SatCom) Networks and focuses on opportunities arising from new technologies and architectures, both in the satellite area and the networking area. The general approach and emphasis is on integrated satellite-terrestrial networks, or even more generally in communications networks that embrace satellite communications, but also potentially other relevant technologies, such as HAPs, and have end-devices of various sizes, capabilities, and mobility characteristics. Small cube satellites (CubeSats) are game changers and expand the networking topologies and capabilities tremendously, but also various problems. With satellite constellations, network topologies can now be rich and the opportunities for a full network in the sky are endless. Combined with networking technologies that can easily differentiate traffic types and their requirements, the new opportunities shine.

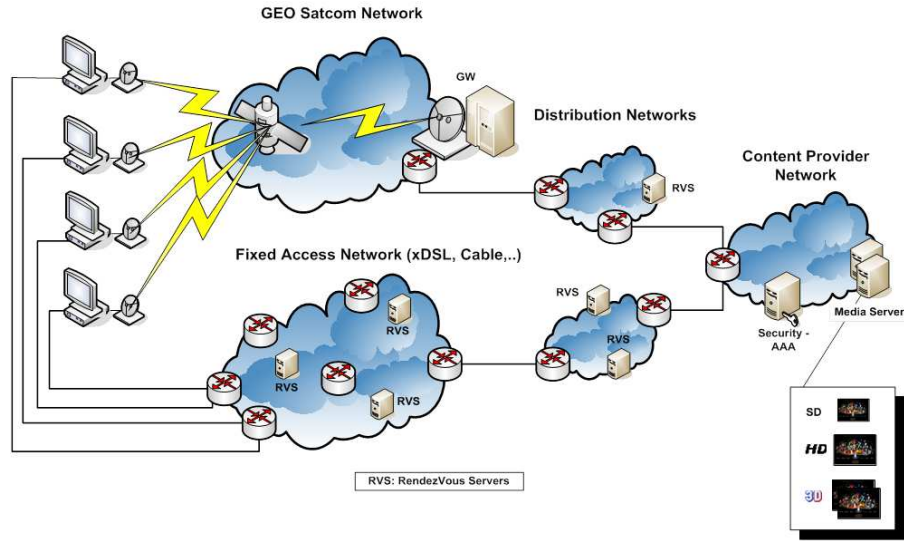


Figure 4.1: IPTV Broadcast.

4.1 SatCom Use Cases and Applications

4.1.1 Content Distribution

IPTV Broadcast

The wide coverage with limited infrastructure needed to be deployed in a geosynchronous (GEO) satellite system continues to make it particularly fitting for TV broadcasting, enabling large economies of scale. In this respect, the current model based on native DVB/MPEG2 transport based broadcasting has significant advantages over terrestrial digital TV, mainly because of the much higher bandwidth available. A hybrid network is another alternative emerging trend. In such a network a SatCom system, based on a typical GEO satellite with classical star topology, is employed together with a terrestrial network for IPTV service (Figure 4.1). The system architecture is depicted in the next figure.

With the respect to the access network, two alternative scenarios can be considered: competitive transmission, and integrated services.

Competitive transmissions imply a form of competition between the SatCom system and the terrestrial one(s), the various configurations depending on the applicable business models. For example, a service provider is customer of two access networks providers: one operating over a satellite network and another operating over a terrestrial networks (e.g, ADSL). End-user contracts with the service provider in this case. The other possibility is that end-user contracts directly with access networks providers while the service management is assured for the access networks by one or several service provider(s) as third-entity. The second alternative, Integrated services, appears as a more promising approach with also an added complexity. What is meant by “integrated services” here results from an association between a terrestrial provider and a satellite provider to offer packaged services with a common management infrastructure and with consistent service delivery (e.g. linear TV programs are all received through the satellite broadcasting and on-demand TV through terrestrial link, etc.). The customer contracts with only one provider that represent all access providers), receives only one bill, and interact with only one common customer service.

VPN in Remote Areas

There cases of services and networks for which a satellite connection is mainly justified from the absence of any possible terrestrial infrastructure. Typical environments for this are offshore platforms, facilities located in remote areas in desert, mountain, or even solution dedicated to remote scientific explorations in areas covered by existing satellite. In this cases a permanent connectivity is required and justifies the cost of a satellite connection of medium or high capacity (at least 1 Mbps on the uplink) to address communication needs. The broadband connection can be used to transport massive amount of data. This requirement may come from the service itself that handles many data (large database distribution for example) or because the remote connection is shared among many other users). Alternately, a real time characteristic may be important to address in some cases (e.g. for distant monitoring from and/or to the isolated location). The typical communication

services may handle very sensitive data for which integrity and/or confidentiality are both essential. In some particular use cases, we can even think that critical data need to be sent or received on time for immediate actions or decisions. Virtual Private Networks (VPN) can be thought of as a convenient way for providing the required isolation with other traffic in the system, both from QoS and security perspective (note that dedicated functions remain of course needed for this). In this use case (Figure 4.2), satellites are typically not used in the first access segment, as users will be connected using a Wi-Fi or LAN Ethernet network. Rather, SatCom systems (based on GEO satellites that are the most natural option here) will feed those access networks as the regional/backbone transit network. What is therefore considered here is a global network (such as the Internet) with multiple access points from which user can be connected to at a given time (sometimes, at the same time). Since transmissions can involve massive data exchanges (in both directions), and because users are supposed to be able to roam from one access point to another one, it appears important that (1) the system is able to deliver service without disruptions (or at least with the minimum of service outage) and (2) in avoiding sending data twice not to waste resource.

At least two modes of transport can be considered: A broadcast /multicast mode (centralized from the satellite gateway or behind) in this case a reliable multicast transport protocol adapted for satellite is needed. A point- to-point mode, if a single location is needed to be fed (e.g. a remote professional site such as an off-shore platform). To this respect, the common PEP may be useless provided that a native transport protocol dedicated to satellite (i.e. taking into account large propagation delays and bandwidth variation if the satellite link is operated in Ka band or above).

Information distribution in disaster areas

Many services such as those needed in Emergency, Public Safety or military communications for mobile and/or personal terminals are considered mission critical. In order to deserve any remote or dangerous areas where terrestrial infrastructures cannot or are not deployed, the access

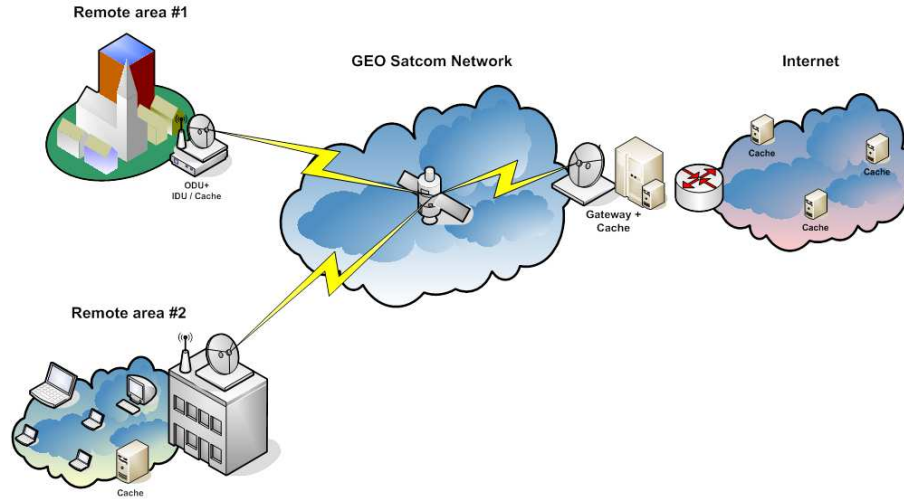


Figure 4.2: VPN in remote areas.

is ensured by a system based on a Low-Earth Orbit (LEO) constellation in very constrained frequency bands (e.g. L-band) compatible with small handheld terminals. Nonetheless, in all other areas the access is provided by a compatible terrestrial technology in order to prevent the use of constraint satellite link and then to increase the system capacity in terms of number of served stations. Satellites and in particular LEOs are the best and most reliable platform for communications in emergency scenarios and perform effectively: when the terrestrial infrastructure is damaged or overloaded, when mobile terminals have to be connected, when a significant part of the communications is only between local actors (e.g. emergency coordination), when broadcast / multicast limited to a small area is required. By nature, a LEO-based system can be assimilated as a “network in the sky” due to the multiple ISLs that constitute a fully mesh network. Such systems are often characterized by path diversity that not only takes place on the space segment (since several satellites may offer concurrent routes) but also on the ground segment because several gateways are deployed in the system and can be selected as entry or exit points to external network

(e.g. the Internet). Mesh capabilities of LEO has another advantage suits to direct user-to-user communications. For users that are relatively close while being isolated (e.g. terminal located under the same spot beam) only a single satellite hop could be sufficient, leading to an optimal routing.

User-Centric Content Sharing

The IST (Information Software Technology) society and the globalization of digital services have shown, in only 10 years that the Internet is far from being a unidirectional network that only delivers anonymously service and data to end-users. Groups or communities have constantly increased from the beginning of the Internet. Communications networks are not only a way to help people to keep in touch; they favour group creation and collaborations. World-wide projects such as Wikipedia or popular open-source software (e.g. some Linux distributions) proves that many people are eager to collaborate, to help each other if they have (even partly) common objective(s). Other example of collaborative networking is computing grids that may apply on small networks (e.g. at the scale of a given organization up to the Internet scale) In the networking domain, the most known collaborative model is P2P communications where clients accepts to become servers in forming a CDN overlay network (Figure 4.3) that increase the data availability and/or quality (Voice and video services over IP). A more recent model, adopted by ISPs targeting residential users, is the “Wi-Fi community” feature where the set-of-box Access Point is configured to offer connectivity to temporary Wi-Fi users that are customers of the same provider if they accept their own box is configured as such.

By definition ad-hoc networks are operated without support of infrastructure (such as fixed access point) and stations may have only limited connectivity (1) with the rest of the ad-hoc network and/or (2) with other networks and the Internet. For wireless ad-hoc networks, most of the time, limitations come from the movement of the node (that enters and exits to/from the area of radio coverage of the ad hoc) in case of wireless communications, the battery limitations, or simply the non-permanent need of service that makes the operator or the user

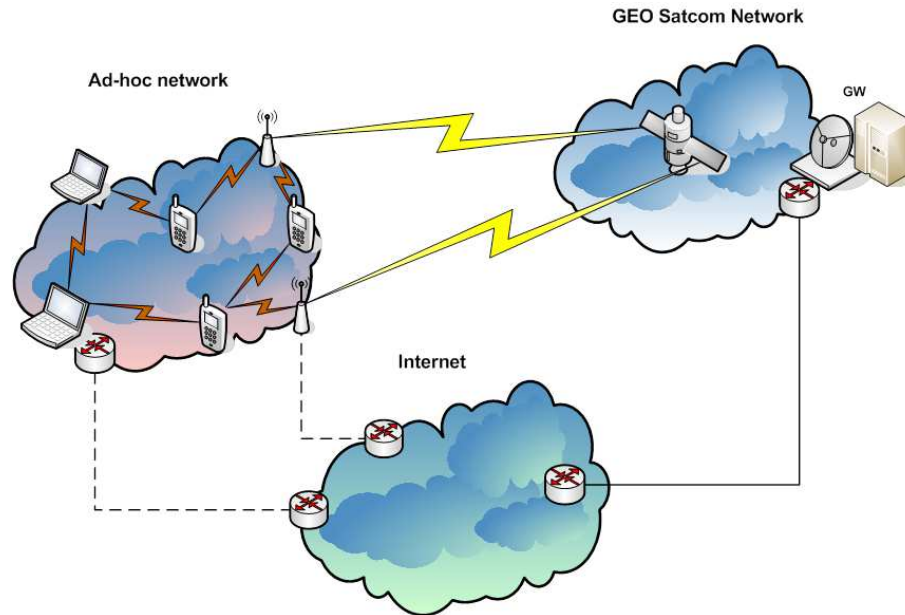


Figure 4.3: User-centric content sharing.

to disconnect or switched off its device(s). On the other hand, satellite provides for those nodes a fixed, permanent and reliable connection that can make the ad-hoc network possibly unlimited in size and in time, provided that the ad-hoc nodes remain in visibility. Several network architectures uses cases can be declined from this scenario, either ad-hoc nodes (at least some of them) have the possibility to establish a satellite connection or there is one or several (fixed) satellite access points that form a minimal and permanent infrastructure in the ad-hoc network. Alternatively, other means of Internet connection may exist in the ad-hoc network via other access points; but this link may not be of use for all services, or it may have intermittent availability with the rest of the ad-hoc network.

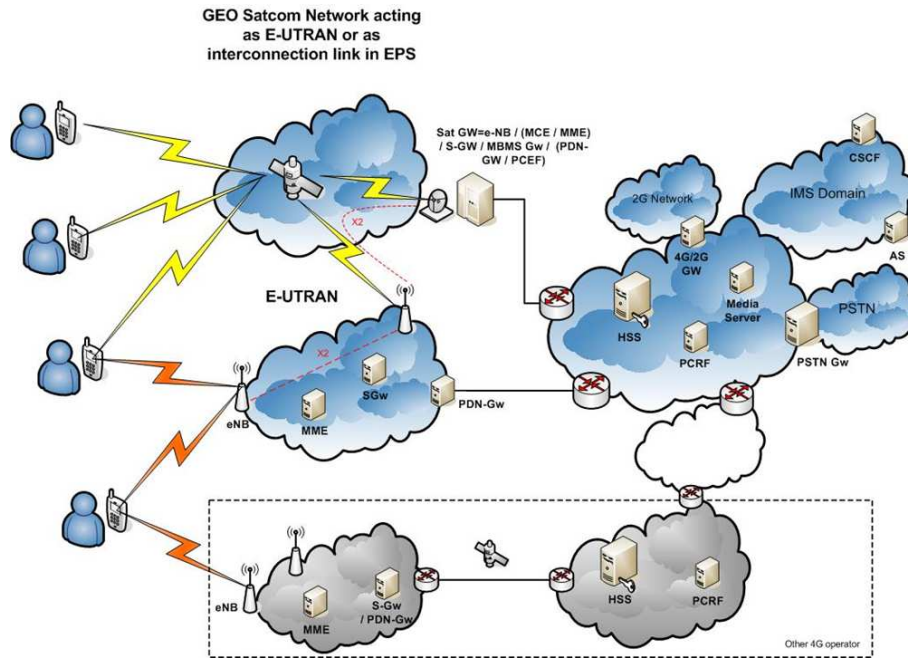


Figure 4.4: 4G-Satellite hybrid content distribution.

4G-Satellite hybrid content distribution

Integrating satellite in a terrestrial 4G infrastructure (Figure 4.4) allows Service Providers to extend their services towards isolated areas. The satellite link is perceived as a means to provide connectivity: when no terrestrial infrastructure / alternative communications channel exists, or its deployment is costly (permanent connectivity), when existing terrestrial infrastructure has been destroyed and service areas must be rapidly covered by an equivalent or even degraded service (temporary connectivity), When it is particularly adapted to a given category of service (specific connectivity). This general context firstly refers to base stations feeds (service usually referred to as backhauling) but also, in some cases, to user terminal direct access via satellite. While these two use cases can duly be considered as different, their technical implementation may share common aspects, in particular due to the flatness

property of 4G architectures. Satellite may be used to limit the remote deployment to eNBs, in areas where it is not technically or economically feasible to deploy a gateway (GWiS-GW / PDN-GW), for example if there are insufficient subscribers (e.g. coverage of small and distant island territory for a national operator). Further, in the model with direct access through satellite, there is no need at all to deploy any e-UTRAN or EPC entity in the remote/difficult areas. This is of importance since the installation of eNBs is costly, has to face environmental/technical constraints and telecom local regulations, all limiting the deployment in space, time, and population coverage. This is particularly true for new sites not equipped with cellular base stations, but this also partially applicable for locations where BS have to be upgraded in eNB. In addition, costs related to multiple eNBs operations and maintenance remain non negligible with respect to a centralized deployment as proposed. In addition, since terrestrial radio spectrum remains a scarce resource taking into account the explosion of mobile traffic expected, it can also be envisaged that satellite resource comes as a complement to terrestrial ones. A general system architecture breakdown is shown on the figure below.

4.1.2 The Internet-of-Things

Vehicle to Vehicle communication

Smart Transport applications based on Machine-to-Machine (M2M) communication are gaining increased popularity as a field suitable to introduce enhanced services and increase efficiency in a number of tangible parameters (e.g. fleet management, fuel consumption, fault prevention, logistics etc.). The fundamental unit of M2M for Smart Transport is the mobile node, which may includes anything from a public or private vehicle (e.g. public urban transport, Taxi vehicles, utility / emergency vehicles, civilian cars etc.) to containers. The M2M aspect of such scenarios involves the exchange of data between vehicles and a centralized platform, or between vehicles themselves. This data could include: location information (i.e. data packets including GPS geographical coordinates) to inform a centralized location monitoring sys-

tem about the exact position of tracked terminals, vehicle related sensor data (e.g. fuel consumption, speed, vehicle condition, engine fault alerts) to provide condition monitoring of vehicles, vehicle unrelated sensor data (e.g. outside temperature measurements, road congestion information, localized hazard information etc.), dispatching information to vehicles (e.g. push a service request to a Taxi or inform a public safety vehicle of an emergency); this could include anything from location information to incident description, operational status information which includes information from-to vehicles regarding the status of e.g. a container delivery or emergency incident resolution, standard communication information (this is not necessarily M2M communications, but falls within the general scope of data exchange between vehicles and centralized platforms). The plethora of information suitable for M2M Smart Transport data communications facilitates a number of applications such as fleet management, public safety applications and traffic congestion management platforms. One of the key design considerations in a M2M Smart Transport scenario involves the definition of the communication paths. Specifically, the following cases are identified: Vehicle-to-Infrastructure (V2I), this is the case when all communication is between vehicles and a centralized Platform (e.g. a fleet management application), Vehicle-to-Vehicle (V2V), this is the case when all communication is between vehicles (e.g. collaborative vehicular applications), Hybrid V2I/V2V, this case combines both aforementioned communication paths, typically, this involves applications that are based on data exchange between both vehicles themselves and a centralized platform. SatCom can provide an efficient framework for M2M Smart Transport applications in a number of scenarios including: cases where terrestrial infrastructure is non-existent or inadequate, in collaboration with terrestrial infrastructure providing a coverage overlay, in M2M Smart Transport applications including frequent broadcasts in vast regions. In providing the framework for M2M Smart Transport, satellite networks can be employed as primary transmission channels or as patch networks in collaboration with terrestrial networks, assuming multiple-interface terminals. The satellite network itself could comprise a GEO, MEO or LEO topology, each resulting in a latency /complexity/mobility sup-

port trade-off. For instance, a GEO network would incur increased delay but results in fewer and less rapid handovers (e.g. beam handover), whereas a LEO constellation provides significantly lower latency but is more susceptible to quick handovers and potentially loss of line of sight. In any case, the satellite network will typically cover V2I communications and potentially V2V communications assuming a star-like topology and non-existent Vehicular Ad-Hoc Network (VANET) links between mobile terminals. In the latter case, LEO constellations are much better suited as a GEO link would result in significant increase in communication delay. By using M2M communications in conjunction with satellite-based GPS units and location-based services companies benefit from real-time information. The data provided can include vehicle location, driver speeds and employee work time. Moreover, developments in communication technology mean it will soon be possible to track goods down to their pick-up or drop-off location, and even confirm delivery and make payments with M2M.

M2M communication between remote objects

Satellite M2M communications, currently a relatively small market compared to that of terrestrial wireless networks, is expanding, led by the growing need to track and monitor difficult-to-reach valuable assets. Satellite M2M networking is gaining popularity in a wide range of governmental and commercial markets owing to the considerable benefits offered by the technology [5]. For specific industries, such as oil and gas, utilities, transportation and logistics, wherein large numbers of assets are deployed in remote locations with inadequate cellular coverage, M2M becomes an expensive proposition unless satellite networks are used (and deployed, if needed). Satellite M2M is preferred for high-end applications requiring wide coverage, particularly in oceans, deserted or rural areas, and for performing e.g., environmental monitoring, mission-critical monitoring, video surveillance etc. Satellite M2M is being revisited and gaining popularity towards the SatCom-enabled 'Internet of Things' paradigm. In the above context of Satellite M2M communications, satellite can play several roles: access segment for part of the devices that can be located outside a terrestrial cover-

age, data aggregation or broadcast for devices spread in a large area, backhauling segment interconnected with few access points that manages aggregated the M2M traffic, reach difficult zones for which no terrestrial infrastructure can be deployed, complement link technology that offer additional capacity for some services / or that may increase system availability. Some devices may be subject to move and the satellite connection can sometimes be temporary, or may provide variable bit-rate connections due to mobility. Two cases can be considered: the satellite connection is deterministic (predictable instant of connection and/or link availability duration) or may be completely random. Also, the traffic model can be very smooth or on the other hand completely burst in response to a local event (e.g. meteorological event or application traffic characteristics).

Smart Grid

One the fastest growing trends in the ICT & Energy convergence is the Smart Grid applications/services framework. While the energy grid, from power generation to power transmission and distribution up to power consumption has been properly functional for decades, it still lacks the automation mechanisms readily available in all other networks to increase its efficiency. The availability of enhanced ICT practices in the area of monitoring (both real and non-real time), analysis and intelligent rules/actions enforcement along with the constantly increasing awareness regarding the necessity of energy efficiency and environmental footprint reduction form the basis on which a number of new applications can be introduced to the energy grid to achieve several goals, such as: Automatic Meter Reading (AMR), early fault detection, preventive maintenance, technical & non-technical loss identification. In general, the objectives of each Smart Grid (Figure 4.5) application fall within one of the following two categories: (a) Collecting information, (b) Enforcing rules based on analysis of collected information. On a second level, Smart Grid applications may refer to one or more of the following parts of the energy grid: power generation, power transmission/distribution, power consumption. While monitoring, control and, ultimately, optimization is desired in each of these 3 operational sets,

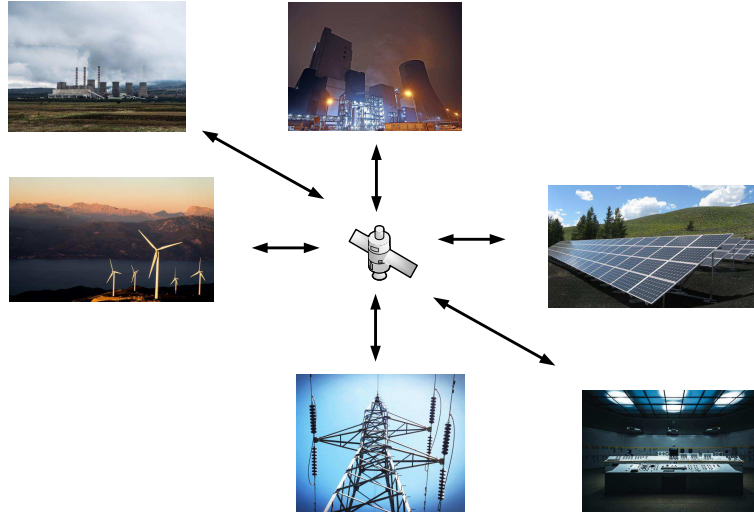


Figure 4.5: Satellite connected smart grid.

as well as any assorted functional subsets, the diversity in infrastructure, business/operational models and involved parties results in multiple roles and access rights with respect to both data collection and actions.

Considering on the one hand the size of the power grid (e.g. overhead/underground power distribution network), the diversity of grid elements (transformers, circuit breakers, inductive loads, meters, conductors etc.) and the potentially immense number of smart grid applications (current/voltage quality measurements, technical loss measurements etc.) it is justified to say that a Smart Grid infrastructure capable of providing heterogeneous services will consist of multiple low power consumption devices each performing a dedicated service (e.g. sensor data), altogether forming ad hoc networks with specific gateway points (e.g. aggregator devices with increased functional capacity). These networks are typically self-organized in an ad-hoc manner. A satellite network can serve as gateway network for the terrestrial ad-hoc infrastructure. From a satellite technology perspective, recent research activities have addressed the issue of low power consumption, random access schemes that are specifically suited for similar M2M

communication applications over satellite. Moreover, data collection from power grid sensors can be improved by using a satellite network serving as the gateway for the terrestrial sensor networks. Collaboration with terrestrial wireless networks (e.g. 3G/LTE) can be addressed, wherein satellite networks manage data collection from areas with limited or no terrestrial connectivity.

4.2 Future SatCom based on ICN architectures

4.2.1 Overview of ICN architectures

The current Internet architecture was founded upon a host-centric communication model, which was appropriate for coping with the needs of the early Internet users. Internet usage has evolved however, with most users mainly interested in accessing (vast amounts of) information, irrespective of its physical location. This paradigm shift in the usage model of the Internet, along with the pressing needs for, among others, better security and mobility support, has led researchers into considering a radical change to the Internet architecture. In this direction, we have witnessed many research efforts investigating Information Centric Networking (ICN) as a foundation upon which the Future Internet can be built [109]. By naming information at the network layer, ICN favors the deployment of in-network caching (or storage, more generally) and multicast mechanisms, thus facilitating the efficient and timely delivery of information to the users. However, there is more to ICN than information distribution, with related research initiatives employing information-awareness as the means for addressing a series of additional limitations in the current Internet architecture, for example, mobility management and security enforcement, so as to fulfill the entire spectrum of Future Internet requirements and objectives.

We now introduce the key concepts and principles of ICN (see [109] for a more extensive tutorial and survey of various ICN architectures).

Key ICN Principles

Focus on Information Identifiers The ICN approach fundamentally decouples information from its sources, by means of a clear location-

identity split. The basic assumption behind this is that information is named, addressed, and matched independently of its location, therefore it may be located anywhere in the network. In ICN, instead of specifying a source-destination host pair for communication, a piece of information itself is named. An indirect implication (and benefit) of moving from the host naming model to the information naming model, is that information retrieval becomes receiver-driven. In contrast to the current Internet where senders have absolute control over the data exchanged, in ICN no data can be received unless it is explicitly requested by the receiver. In ICN, after a request is sent, the network is responsible for locating the best source that can provide the desired information. Routing of information requests thus seeks to find the best source for the information, based on a location-independent name (or identifier).

Focus on Information Delivery In ICN the network may satisfy an information request not only through locating the original information source, but also by utilizing (possibly multiple) in-network caches that hold copies of the desired information (or pieces of it). This can be accomplished without resorting to add-on, proprietary and costly overlay solutions (e.g., Content Delivery Networks-CDNs), since the network layer in ICN operates directly on named information. ICN-based architectures see non-opaque data packets, in the sense that these are named based on the information they carry. Therefore, information fragments (packets in current terms) can be cached and retrieved easily, unlike in the current Internet where costly techniques like Deep Packet Inspection (DPI) would have to be used for answering requests with data/packets cached on the routers, not to mention that DPI does not work when packets are encrypted. Moreover, by naming information, ICN allows the aggregation of requests for the same information, thus facilitating its delivery to the corresponding destinations via multicast forwarding. Finally, access control (i.e., who is allowed to access which data) can potentially be applied directly at the network layer, since network elements are aware of what information is being transferred inside each packet.

Focus on Mobility In ICN, host mobility is addressed by employing the *publish/subscribe* communication model. In this model, users interested in information subscribe to it, i.e., they denote their interest for it to the network, and users offering information publish advertisements for information to the network. Inside the network, brokers are responsible for matching subscriptions with publications i.e., they provide a rendezvous function. It is important to note that the publish/subscribe terminology used in the context of ICN differs from that of traditional (application layer) publish/subscribe systems. In traditional publish/subscribe systems, publish involves the actual transmission of data while subscribe results in receiving data published in the future, with the ability of receiving previously published data being optional. In ICN, on the other hand, publish involves only announcement of the availability of information to the network, whereas subscriptions by default refer to already available information (or perhaps information items to become available in the future but to be served once), leaving the option of permanent subscriptions (i.e., receiving multiple publications matching a single subscription) as optional. The strength of the publish/subscribe communication model stems from the fact that publication and subscription operations are decoupled in time and space. The communication between a publisher and a subscriber does not need to be time-synchronized, i.e., the publisher may publish information before any subscribers have requested it and the subscribers may initiate information requests after or before publication announcements. Publishers do not usually hold references to the subscribers, neither do they know how many subscribers are receiving a particular publication and, similarly, subscribers do not usually hold references to the publishers, neither do they know how many publishers are providing the information. These properties allow for the efficient support of mobility: mobile nodes can simply reissue subscriptions for information after handoffs and the network may direct these subscriptions to nearby caches rather than to the original publisher.

Focus on Security Many of the security problems of the Internet are largely due to the disconnection between information semantics at the

application layer and the opaque data in individual IP packets. This places a significant burden on integrating accountability mechanisms into the overall architecture. Point solutions like DPI or lawful interception try to restore this broken link between the actual information semantics and the data scattered in individual packets. However, this is achieved at a relatively high cost and is therefore only applicable to critical problems, such as law enforcement. As a result, while secure end-to-end connections are prevalent, the overall Internet architecture is still not self-protected against malicious attacks and data is not secure. At the same time, the lack of an accountability framework which would allow non-intrusive and non-discriminatory means to detect misbehavior and mitigate its effects, while retaining the broad accessibility to the Internet and ensuring both data security and communication privacy (i.e., hiding from non-authorized parties that a communication between two parties took place) is a crucial limitation to overcome. ICN architectures are in contrast interest-driven, i.e., there is no data flow unless a user has explicitly asked for a particular piece of information. This is expected to significantly reduce the amount of unwanted data transfers (such as spam) and also facilitate the deployment of accountability and forensic mechanisms on the network points that handle “availability” and “interest” signaling. Moreover, for ICN architectures that use self-certifying names for information, malicious data filtering will be possible even by in-network mechanisms. Finally, most ICN architectures add a point of indirection between users requesting a piece of information and users possessing this piece of information, decoupling the communication between these parties. This decoupling can be a step towards fighting denial of service attacks, as requests can be evaluated at the indirection point, prior to arriving to their final destination. Indirection can also benefit user privacy, as a publisher does not need to be aware of the identities of its subscribers.

4.2.2 ICN and Integrated Satellite-Terrestrial Networks

In this section we discuss how integrated satellite-terrestrial networks can benefit from ICN. In particular, we present a specific ICN archi-

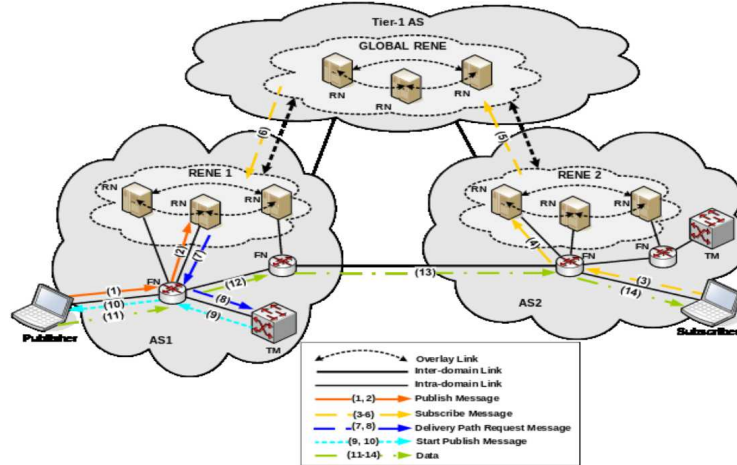


Figure 4.6: The PSI architecture. RN stands for Rendezvous Node, RENE for RENEzvous Network, FN for Forwarding Node and TM for Topology Manager. (This figure is from [109].).

ecture, namely Publish-Subscribe Internetworking (PSI)¹, and an integration plan and we examine advantages ICN can bring in supporting SatCom applications and use cases. The overall architecture is sketched in Figure 4.6).

Publish-Subscribe Internetworking (PSI)

PSI is a clean-slate Internet architecture that consists of three key separate functions supporting information exchange through the publish-subscribe operation: rendezvous, topology management and forwarding. When the rendezvous function matches a subscription to a publication, it directs the topology management function to create a route between the publisher and the subscriber. This route is finally used by the forwarding function to perform the actual transfer of data.

¹A more extensive (but also tutorial style) description of the PSI architecture can be found in [108]

Naming Information objects in PSI are identified by a (statistically²) unique pair of IDs, the scope ID and the rendezvous ID. The scope ID groups related information objects while the rendezvous ID is the actual identity for a particular piece of information. Information objects may belong to multiple scopes (possibly with different rendezvous IDs), but they must always belong to at least one scope. Scopes serve as a means of (a) defining sets of information objects within a given context and (b) enforcing “boundaries” based on some dissemination strategy for the scope. For example, a publisher may place a photograph under a “friends” scope and a “family” scope, with each scope having different access rights. While PSI rendezvous IDs (names) are flat, scopes can be organized in scope graphs of various forms, including hierarchies, therefore a complete name consists of a sequence of scope IDs and a single rendezvous ID.

Name Resolution and Data Routing Name resolution is handled by the rendezvous function, which is implemented by a collection of Rendezvous Nodess (RNs), the Rendezvous Network (RENE), as shown below. The RENEs could be implemented as Distributed Hash Tables (DHT) and the global RENE as a hierarchical DHT [41]. When a publisher wants to advertise an information object, it issues a PUBLISH message to its local RN which is routed by the DHT to the RN assigned with the corresponding scope ID (arrows 1, 2). When a subscriber issues a SUBSCRIBE message for the same information object to its local RN, it is routed by the DHT to the same RN (arrows 3 to 6). The RN then instructs the Topology Manager (TM)) to create a route connecting the publisher with the subscriber for data delivery (arrows 7, 8). The TM sends that route to the publisher in a START PUBLISH message (arrows 9, 10) that finally uses this route to send the information object via a set of Forwarding Nodes (FNs). The TM implements the topology management function through a distributed routing protocol

²This choice is made so that rendezvous IDs can be chosen by publishers independently. That the same ID refers to different information items can be determined e.g., by the rendezvous node without having access to the actual item by considering the metadata. If a ‘collision’ in rendezvous IDs is detected, procedures might be in place to resolve the collision, e.g., by requesting the change of one of the IDs.

(in order to discover the network topology), e.g., OSPF. The actual delivery paths are calculated upon request by the rendezvous function as a series of links between FNs and encoded into source routes using a technique based on Bloom filters. Specifically, each network node assigns a tag, i.e., a long bit string produced by a set of hash functions, to each of its outgoing links, and advertises these tags via the routing protocol. A path through the network is then encoded by *ORing* the tags of its constituent links and the resulting Bloom filter is included in each data packet. When a data packet arrives at a FN, the FN simply *ANDs* the tags of its outgoing links with the Bloom filter in the packet; for every tag that matches, the packet is forwarded over the corresponding link. In this manner, the only state maintained at the FNs is the list of link tags. Multicast transmission can be achieved by simply encoding the entire multicast tree into a single Bloom filter.

Subsequent packets belonging to the same information object can be individually requested by the subscriber using e.g. the notion of Algorithmic IDs, i.e., packet IDs generated by an algorithm agreed by the communicating entities. These requests are forwarded similarly to data packets; by using reverse Bloom filters calculated by the TM they could bypass the RENE. This allows the realization of transport layer protocols, e.g., via a sliding window of pending requests. Furthermore, this can be generalized to multi-source and multipath transport [98].

Name resolution and data routing are decoupled in PSI, since name resolution is performed by the RENE, while data routing is organized by the TMs and executed by the FNs. While name resolution can be time consuming, data forwarding can take place at line speeds, without placing any state at the FNs [50]. Furthermore, the separation of routing and forwarding allows the TMs to calculate paths using complex criteria (e.g., load balancing), without requiring signaling to the (stateless) FNs. On the other hand, the topology management and forwarding functions as described here are only adequate for the intra-domain case and need to be extended (e.g., with label switching) for the inter-domain level.

Caching PSI can support both on-path and off-path caching. In the on-path case, forwarded packets are cached at FNs in order to potentially serve subsequent requests. However, on-path caching may not be very effective due to the decoupled nature of name resolution and data routing, since requests for the same information object can reach the same RN, whereas the actual data transfers may use entirely different paths. In the off-path case, caches operate as publishers, by advertising the available information to the RENE. Managed information replication (in 'surrogates'), as in CDNs, can also be efficiently supported by the PSI architecture [108].

Mobility Mobility is greatly facilitated by the use of multicast and caching. Four types of mobility cases are considered, based on host movement (local, global) and technology (static when a single technology is involved, dynamic when vertical handoffs are performed). Local subscriber mobility can be handled via multicast and caching, i.e., by multicasting information objects to the multiple possible locations for the mobile subscriber and the mobile subscriber receiving information objects from nearby caches after a handoff. Global subscriber mobility is handled by modifying the forwarding function of the architecture. Mobility prediction can be used to reduce handoff latencies by caching information requested by the subscriber to the areas where the subscriber is expected to move after a handoff [99]. Publisher mobility is harder, since the topology management function has to be notified of the publisher's new position in the network (but could be easily bypassed by mobiles having anchor points (i.e. proxies) in the fixed network).

Security The publish-subscribe paradigm can be seen as a remedy to the imbalance of power between senders and receivers in the traditional send-receive paradigm. This imbalance is often accused for the increasing number of (Distributed) Denial of Service (DDoS) attacks, as well as for the emergence of spamming. In pub/sub systems there is no information flow as long as the receiver has not expressed interest in a particular piece of information, i.e, the receiver in a pub/sub architec-

ture is able to instruct the network which pieces of information shall be delivered to it. Paths encoded into Bloom filters can use dynamic link identifiers, making it impossible for an attacker to craft Bloom filters or even to reuse old Bloom filters to launch DoS attacks [Jok2009].

Moreover, and even though the model is so powerful so that there can be subscriptions before the corresponding publications have been published, no information is requested from a publisher, unless the publisher has explicitly denoted the availability of that information, i.e., not before the publisher has issued a publication message (for this particular piece of information). Publication and subscription operations are decoupled in time and space, i.e., they do not have to be synchronized neither do they block each other. Moreover publishers and subscribers do not communicate directly and they can hide their identity as, in general, subscribers are only interested in the information itself rather than on who provides it, and publishers, usually, disseminate publications using multicast so they cannot (and usually should not) be fully aware of the publication's recipients. Therefore, anonymity can be easily achieved in pub/sub architectures. PSI also supports Packet Level Authentication (PLA) technique for encrypting and signing individual packets. This technique assures data integrity and confidentiality as well as malicious publisher accountability. PLA can be used to check packets either at FNs or at their final destination. The use of flat names also permits self-certifying names for immutable data objects, using the object's hash as the rendezvous ID and checking the received item against the self-certifying name. Additionally, with the RN a point in the network where subscriptions and publications are matched (and with metadata being available) effective deployment of access control mechanisms is enabled. Furthermore, lightweight, fast application of access control to the ubiquitous storage in the network is necessary and possible, e.g. through access-control delegation [42].

Publish-Subscribe architectures also offer good availability because the rendezvous network is usually implemented using a DHT and DHTs provide effective load balancing—usually at the cost of some communication stretch. Moreover multi-homing can be easily supported, as multiple publishers may advertise the same publication to a RENE

(or even multiple independent, parallel RENEs), therefore a RN has a number of options with which it can satisfy a subscription, improving availability. Furthermore, subscription aggregation can easily be realized and exploited through multicast improving resource sharing, which leads again to greater availability. Finally, security solutions for mitigating spam and for preserving privacy have been developed for PSI.

Migrating integrated satellite-terrestrial networks to PSI

In this section we discuss how core PSI functionalities can be extended and adapted for SatCom.

Name Resolution Due to the potentially large amount of memory that a RN may require, implementing the rendezvous functionality on satellites will not, in most cases, be a feasible solution. This is typically true for bent-pipe satellite systems (transparent GEO satellites). However there may be real benefits for having such rendezvous function on-board in particular missions or configurations:

- Signaling traffic may be significantly reduced on the feeder link, if rendezvous exchanges initiated from satellite terminals would not have to be forwarded to the ground, to reach the gateway.
- At the same time, rendezvous delays may also be significantly decreased. This may be important for low delay services, when published data must be received by subscribers as soon as possible.

From a practical point of view, one could only consider RNs in LEO constellations systems and for a small amount of urgent data. This can only address very specific missions. Note that this particular deployment assumes the presence of a secondary, terrestrial, connection.

Data Routing Bloom filter based data forwarding can be used in the satellite environment, with the forwarding trees within the satellite domain ending (or starting) at GWs. This way the GWs act as the

dedicated forwarders for satellite terminals (be they subscribers or publishers). This means that in a sense a GW keeps track of which nodes use it as a proxy for publishing or subscribing to data. For that reason, in case the publishers/subscribers are not satellite terminals but separate end-hosts, it could be reasonable to setup a proxying function to increase scalability at the GW level. This way the GW acts as a bridge where on the network side it forwards packets using Bloom filter based forwarding, while on the end host side it can use any forwarding mechanism by using a table that maps the forwarding identifiers to end host addresses and vice versa. This way even if the reason for directing traffic towards a GW is due to a false positive³, the GW can drop it if according to its mapping tables no-one has registered for it (and hence the satellite links will not be used and remain unaffected by this false positive). Also, to avoid false positive routing at GW with multi-satellite links (mesh systems, multi-spots, multi-satellite systems, i.e. systems with several carriers transmitted from the GW) some extra checks can be performed on top of the normal Bloom filter based forwarding operations to ensure that no unnecessary link will be used.

Topology Management One local TM can be used for the satellite domain. For multi-GW systems, each GW could host a local TM that maintains the view of the whole satellite network, based on link information it receives from satellites connected to it. The view can consist of the network topology and the load of various links. In case of a distributed implementation the TMs may share reports among them (preferably via terrestrial links) such that any TM can respond to queries sent by the RENE on the best paths to interconnect publishers with subscribers. In a centralized approach all local TMs may send their reports to one central (domain) TM, which would be responsible for answering queries coming from an RN on the best paths within the satellite network.

³False positives may occur because of the use of Bloom filters and result in packets being (mis)routed over additional paths in the network; they can be controlled by limiting the diameter of the network [50]

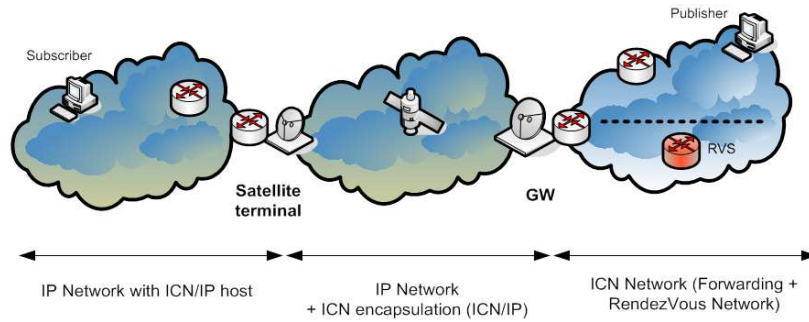


Figure 4.7: ICN integration (short-term).

Deployment and transition issues

The question of the target time horizon for deploying a new architecture and transition issues cannot be avoided. A main concern here is the long development cycles that prevail in the satellite industry. This matter primarily concerns the space segment, however the ground segment (gateway and terminals) is also impacted. In practice it may take many years to have systems evolve. The following migration strategies can be considered.

Short-term migration Tunneling solutions over existing technologies can be developed, such as “*ICN over IP*” (or in alternate forms ICN over TCP or ICN over UDP). The impact on the core satellite network implementation will be low, as existing interfaces and control-plane functions will not be required to change. To minimize change with respect to the satellite, the RENE architecture could be completely split and independent from the SatCom network and could be managed by a separate service provider. The overall concept is illustrated in Figure 4.7.

Mid-term migration In the longer-term, a migration to ICN could be realized with the satellite access network still based on existing forwarding interfaces, but also having to interconnect native ICN networks

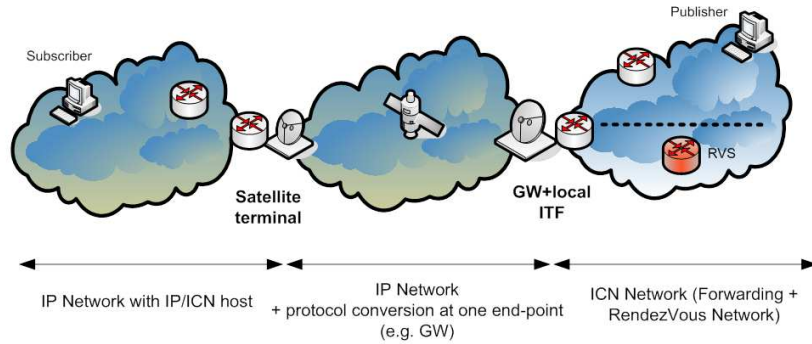


Figure 4.8: ICN integration (mid-term).

that would not support IP. In that case, there must be at least a local inter-/intra-domain Topology Management function hosted at the satellite segment. An alternative could be to outsource this function to the ICN transport network interconnected to the satellite domain with knowledge of the satellite topology, link status, available resources etc. The overall concept is illustrated in Figure 4.8.

Probably later, a local rendezvous function could be supported in addition by the SatCom system at least within a centralized architecture. The need to support rendezvous at the GW is mainly tied with the system mission. For example, for a generic broadband access system this might not be a critical issue. On the other hand, for any other systems providing many/important value added services (such as broad content distribution), implementing RNs within the satellite domain would be preferable for better flexibility in the control of information and content management (scoping, filtering, and aggregation).

Long-term migration In the long-term, the satellite terminals and gateways can be either completely dual-stacked (IP/ICN) or implement only ICN interfaces and functions. The satellite networking specificities are implemented through customized policies for all core functions of the PSI architecture. The satellite gateways (or dedicated servers in the vicinity of the gateway managed by the SatCom access provider)

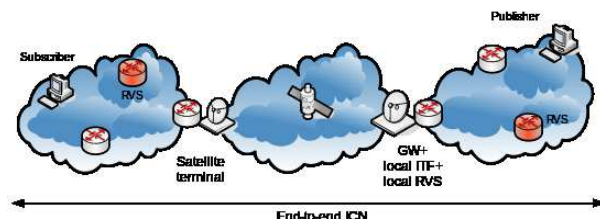


Figure 4.9: ICN integration (long-term).

would typically implement an RN, a TM, and caching. The GW would now act as an ICN forwarding node. The overall concept is illustrated in Figure 4.9.

Impact of ICN on Integrated Satellite-Terrestrial Networks

We now examine how ICN affects the following scenarios: (a) hybrid broadcast IPTV, (b) smart M2M transport, (c) extended cellular (4G) backhauling. For each of these scenarios, the associated issues are presented and the PSI relevance and benefits are elaborated. The hybrid broadcast IPTV scenario is discussed in detail, while the subsections for the other two scenarios describe the main differences with respect to the first one.

Hybrid Broadcast IPTV Scenario This scenario paves the way for SatCom integration with the Future Media Internet. Typical use case is a GEO SatCom system with the classical star topology integrated with a terrestrial network. Two or three separate actors are involved in the service provision, namely Content Providers (CPs), Access Providers (APs) (satellite and terrestrial) and optionally Transit Providers (TPs) to interconnect content and access providers. The overall architecture is shown in Figure 4.10.

The PSI architecture favors collaborative service provisioning where the CPs “poll” the access networks to determine the optimal forward path using the hierarchical organization of PSI functions with the CP at the top level. CPs could manage the core RENE and servers. This network could also host the Inter-Domain Topology Manager or at

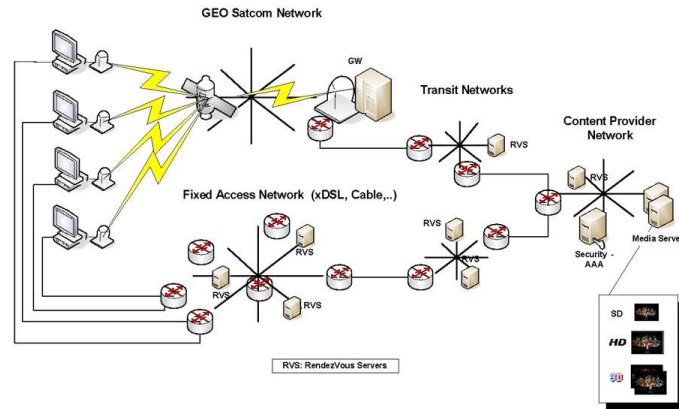


Figure 4.10: Hybrid broadcast IPTV scenario.

least one local Topology Manager (i.e. if one TM is allocated to each access domain). A local RN can also be implemented near the satellite gateway for load balancing purposes. This design choice allows CPs to manage the dissemination of content within their own networks and assist access networks to localize published content as near as possible to forwarding nodes. Levels of collaborative service provisioning depend on both economic (i.e. cost minimization) as well as technical criteria (i.e. QoS guarantees, required bandwidth, date and time of planned broadcast, number and geographical distribution of subscribers to specific content, popularity of content, etc.) and may lead to:

- Influencing routing decisions (i.e. which network to use).
- Supporting the implementation of switching from one connection mode to another one (e.g., in case of link failure).
- Supporting load-balancing between networks.
- Splitting service flows into several components sent separately over the two networks (the strategic decision being in that case how the splitting is done); e.g., employing layered video coding, such as Scalable Video Coding (SVC) or Multi-View Video Coding (MVC), send the base layers through the satellite and the enhancement layers through terrestrial.

- Sending traffic from one network and receive it from the other one; particularly attractive for asymmetrical satellite systems.
- Sharing RNs between networks.
- Sharing caches (increasing the caching capacity as seen by each network), particularly relevant for on demand services as content can be dynamically cached near the satellite gateway and the terrestrial access point

Routing protocols like Multi-Protocol Label Switching (MPLS) (TE, with traffic engineering extensions) or their evolutions in the context of the Future Internet, or simpler mechanisms like Policy-Based Routing (PBR) rules (e.g. as supported by Cisco IOS-based routers), are examples of implementations for the QoS-related forward path and network selection. Other important technical aspects to optimize the performance of PSI-based SatCom within the setting of on-demand services are cache management (which content items to cache and where) and timing policies to govern the update and data overwriting. Decisions on cache selection may be done jointly with the selection of the access network.

Smart M2M Transport The scenario paves the way for SatCom integration with Internet of Things (IoT). SatCom can provide an efficient framework for M2M smart transport services in maritime, aeronautical and fast rail environments, in particular:

- When terrestrial infrastructure is non-existent or inadequate.
- In collaboration with terrestrial infrastructure, providing coverage extension.
- When enhancing offered services, i.e. with satellite-based GPS units and location-based services end users benefit from real-time information (e.g., location based data delivered on vehicles).

The satellite network itself could be on GEO, MEO or LEO satellites and would typically cover V2I communications and potentially

also support V2V communications extending VANETs when vehicles move away beyond V2V range. By nature, a LEO-based system can be assimilated as a “network in the sky” due to the multiple ISLs that constitute a fully mesh network. Mesh LEO communications with a single satellite hop could lead to an optimal routing. In this case, static strategies within the constellation are possible since the topology is purely deterministic, but it can also be expected that actual dynamic routing based on on-board processing and taking QoS as well as link loads into account will further be developed in the future. In this way, a good short-term option is based on pure IP routing, as demonstrated by the recent Cisco IP Routing in Space (IRIS) initiative[27]. In a mid-term horizon, satellite storage capacities are expected to increase and this could be instrumental for a large number of services. More generally, on-board capabilities of LEO satellites represent good opportunities to process and/or store data, as required by PSI features. LEO networks could not only act as (one of the) access networks but also host ICN functions, in particular, some name resolution (rendezvous) and routing functions (TM) could be directly implemented on LEO satellites. A first step in the scenario would be that all PSI-related metadata are stored on-board which would only require limited storage and processing resources for the satellites. This would result in a trade-off between PSI performance and processing load in LEO nodes (consequently this could affect applicability from a techno-economical perspective). In case more storage resources were available, but also with large link capacities, on-board caching would bring added QoS by reducing delay and possibly enhancing throughput. Other content-oriented functions (filtering, aggregation, etc.) could also be advantageously integrated on-board. Lastly, LEO constellations result in significant decrease in propagation delay, which also directly impacts the global Quality of Experience (QoE). Important technical aspects to consider in order to optimize the performance of PSI-based SatCom within the context of smart M2M transport services are: (a) hierarchical naming, (b) opportunistic content forwarding and caching, and (c) publishers’ mobility. M2M data are collected from multiple mobile terminals, falling under different and varying scopes with heterogeneous metadata. The metadata are spatial

and temporal with some timestamp based validity period. An efficient naming hierarchy must be used to facilitate the following: information object classification; information prioritization (e.g. alerts such as safety or vehicle related conditions vs. typical measurements); scoping etc. Since many M2M applications include the transfer of multiple data items from remote sources to service centers, handling all such data as a separate publication could cause scalability problems. Naming techniques are being investigated with respect to their ability to support information aggregation. Typically, this would also involve aggregation at higher layers (e.g. the rendezvous and topology management functions). To maximize the energy efficiency, opportunistic content push cache management techniques could be used. There are three main aspects of content storing that need to be addressed depending on the case: (a) caching in SatCom network nodes (e.g. satellite gateway or LEO nodes); (b) opportunistic content forwarding to terrestrial gateways in case of SatCom-terrestrial network integration; (c) caching in mobile terminals; the latter perceived as an enhancement to the proposed VANET architectures. Specifically, a vehicle within the context of VANET could store content to relay to another vehicle (through the terrestrial network). However, this function can be expanded to effectively make vehicles PSI nodes, i.e. nodes that can satisfy requests for content via cached copies.

In M2M scenarios, both subscribers and publishers can be mobile. To reduce the convergence time expected for the RENE update in case of publisher mobility, centralized rendezvous functions can be considered. Implications could arise due to publisher mobility between satellite and terrestrial networks. Finally, on top of the generic PSI benefits, important additional benefits emerge from the inherent support of Delay/Disruption Tolerant Network (DTN)); content originating from or directed to vehicles may be subject to unexpected delays due to several causes, such as temporary loss of line-of-sight (LOS) due to blockage in an urban environment, or by weather conditions impacting SatCom.

Extended 4G Backhauling Scenario This scenario provides advantages in SatCom integration within the cellular/mobile Internet context given the aggressive video traffic growth levels in (terrestrial) mobile networks. Integrating SatCom into a terrestrial cellular infrastructure allows Service Providers to extend their services towards isolated areas and enhance network capabilities for more efficient video content delivery, exploiting e.g. Evolved Multimedia Broadcast Multicast (eMBMS). The satellite network itself could comprise typically a GEO or MEO topology. In some areas terrestrial repeaters (such as in systems based on DVB-SH and -NGH) could also be used to boost the satellite signal, e.g. for reception in urban areas (a hybrid satellite-LTE system).

Major PSI integration benefits arise from the QoS based inter-component handover management. Inter-component handovers in Long-Term Evolution (LTE), i.e. handovers between LTE Radio Access Network (RAN) nodes with different backhaul technologies, lead to performance degradation due to the high variation in delay. Typically, the PSI TM function, which would handle a case of inter-component handover, can include QoS mechanisms in the path selection and data forwarding. Moreover, proactive or opportunistic forwarding and content caching support could assist further seamless handovers by forwarding and caching content at the base stations with satellite backhaul to prevent any noticeable to the mobile user connection changes in terms of delay variation. Seamless service continuity can be realized by transmissions of cached content during the initial period after handover. On the other hand, delays occur also due to the handover signaling over the satellite channel. This can also be mitigated in streaming applications (i.e. on-going data delivery to mobile users) via cached content. Similar benefits of PSI can be perceived in the case of handovers from a satellite backhauled base station to another base station, again via the use of caching. In this case, content delivery is unaffected from the RTT change (decrease) that could cause content to be delivered out of sequence, as the initial phase after the handover can be addressed by transmissions of cached content from the base stations. As previously mentioned, joint optimization of PSI core functions with QoS support functions, satellite resource utilization and bandwidth-

on-demand mechanisms would further optimize integrated networks performance. For example, with respect to traditional backhauling services, the network operators may control finely the delivery, caching and management of content as desired, instead of relying on fixed capacity allocations where some resources would unavoidably be underutilized when the amount of traffic to transport is reduced. Individual access control and billing can also be applied, i.e., per object, per user, or per user and terminal (online or offline charging, object price etc.), or even the decision to route objects via the terrestrial network.

Finally, utilizing the satellite medium as a backhaul link deviates from acceptable security policies regarding core network security in cellular networks. However, within the context of PSI networking, both the content itself and the subscriber in the pub/sub paradigm are authenticated, easing such concerns.

4.2.3 Conclusions

Integrated Satellite-Terrestrial networks have a great potential in the context of future Information-Centric or Content-Centric Internet, allowing great flexibility of providing users and providers the flexibility, performance, security and robustness sought.

References

- [1] Irtf network coding research group (nwcrg) web site with url [date of access may 15]. <https://irtf.org/nwcrg>.
- [2] Leosat enterprises contracts first customer. <http://tinyurl.com/gln7vsp>.
- [3] Web site of the multipath tcp linux kernel implementation. <http://multipath-tcp.org>.
- [4] Wikipedia page of multipath tcp with url [date of access may 15, 2016]. https://en.wikipedia.org/wiki/Multipath_TCP.
- [5] SAMOS (Satellite Machine-to-Machine Services Market Survey), Final Report, 1998.
- [6] B. Adamson, C. Bormann, M. Handley, and J. Macker. Nack-oriented reliable multicast (norm) transport protocol. Technical report, IETF RFC 5740, November 2009.
- [7] R. Ahlswede, N. Cai, S. y. R. Li, and R. W. Yeung. Network information flow. *IEEE Transactions on Information Theory*, 46(4):1204–1216, July 2000.
- [8] I. F. Akyildiz, Xin Zhang, and Jian Fang. Tcp-peach+: enhancement of tcp-peach for satellite ip networks. *IEEE Communications Letters*, 6(7):303–305, July 2002.
- [9] X. Alberti et al. System capacity optimization in time and frequency for multibeam multi-media satellite systems. In *Advanced Satellite Multimedia Systems Conf.*, pages 226–233, Sep. 2010.

- [10] P. Angeletti, R. De Gaudenzi, E. Re, and N. Jeannin. Multibeam satellite communication system and method, and satellite payload for carrying out such a method. *U.S. Patent WO 2014 001 837*, 1, Jan. 2014.
- [11] P. Angeletti, C. Mangenot, and G. Toso. Recent advances on array antennas for multibeam space applications. In *2013 IEEE Antennas and Propagation Society International Symposium (APSURSI)*, pages 2233–2234, July 2013.
- [12] J. Anzalchi and M. Harverson. Generic flexible payload technology for enhancing in-orbit satellite payload flexibility. In *25th AIAA International Communications Satellite Systems Conference*, Sept. 2007.
- [13] A.I. Aravanis, B. Shankar, M.R., P.D. Arapoglou, G. Danoy, P.G. Cottis, and B. Ottersten. Power allocation in multibeam satellite systems: A two-stage multi-objective optimization. *IEEE Trans. on Wireless Commun.*, 14(6):3171–3182, June 2015.
- [14] T. Azzarelli. Oneweb global access. <http://tinyurl.com/z4jtcwl>.
- [15] S. Barré, C. Paasch, and O. Bonaventure. *MultiPath TCP: From Theory to Practice*. IFIP Networking, Valencia, 2011.
- [16] L. Bertaux, S. Medjiah, P. Berthou, S. Abdellatif, A. Hakiri, P. Gelard, F. Planchou, and M. Bruyere. Software defined networking and virtualization for broadband satellite networks. *IEEE Communications Magazine*, 53(3):54–60, March 2015.
- [17] J. Border, M. Kojo, J. Griner, G. Montenegro, and Z. Shelby. Performance Enhancing Proxies Intended to Mitigate Link-Related Degradations. RFC 3135 (Informational), June 2001. <http://www.ietf.org/rfc/rfc3135.txt>.
- [18] M. Bousquet, J. Radzik, M. Jeannin, L. Castanet, P. Thompson, and B. G. Evans. Broadband access terabit/s satellite concepts. In *Proc. of the 29th AIAA International Communications Satellite Systems Conference*, 28 November - 1 Nara, Japan, December 2011. ICSSC.
- [19] Tom Butash and Barry Evans. Ijscn special issue on ka band and high throughput satellites. *International Journal of Satellite Communications and Networking*, 34(4):459–460, 2016.
- [20] M. A. Vazquez Castro and G. S. Granados. Cross-layer packet scheduler design of a multibeam broadband satellite system with adaptive coding and modulation. *IEEE Trans. on Wireless Commun.*, 6(1):248–258, Jan 2007.

- [21] A. Catalani, L. Russo, O. M. Bucci, T. Isernia, G. Toso, and P. Angeletti. Ka-band active sparse arrays for satcom applications. In *Proceedings of the 2012 IEEE International Symposium on Antennas and Propagation*, pages 1–2, July 2012.
- [22] Shigang Chen and K. Nahrstedt. An overview of quality of service routing for next-generation high-speed networks: problems and solutions. *IEEE Network*, 12(6):64–79, Nov 1998.
- [23] J. P. Choi and V. W. S. Chan. Optimum power and beam allocation based on traffic demands and channel conditions over satellite downlinks. *IEEE Trans. on Wireless Commun.*, 4(6):2983–2993, Nov. 2005.
- [24] J. P. Choi and V. W. S. Chan. Resource management for advanced transmission antenna satellites. *IEEE Trans. on Wireless Commun.*, 8(3):1308–1321, Mar. 2009.
- [25] P. A. Chou and Y. Wu. Network coding for the internet and wireless networks. Technical Report Microsoft Research, MSR-TR-2007-70, June 2007.
- [26] G. Cocco, T. D. Cola, M. Angelone, and Z. Katona. Radio resource management strategies for dvb-s2 systems operated with flexible satellite payloads. In *2016 8th Advanced Satellite Multimedia Systems Conference and the 14th Signal Processing for Space Communications Workshop (ASMS/SPSC)*, pages 1–8, Sept 2016.
- [27] E. Cuevas, H. Esiely-Barrera, H. Warren Kim, and Z. Tang. Assessment of the internet protocol routing in space - joint capability technical demonstration. *Johns Hopkins APL Technical Digest*, 30(2):89–102, 2011.
- [28] Peter B. de Selding. Leosat corporate broadband constellation sees geo satellite operators as partners. <http://tinyurl.com/z8o8shg>.
- [29] J. Detchart, E. Lochin, J. Lacan, and V. Roca. Tetrys, an on-the-fly network coding protocol. Technical report, NWCRG Internet Draft, July 8, 2016.
- [30] M. Durvy, C. Fragouli, and P. Thiran. Towards reliable broadcasting using acks. In *Proc. IEEE Int. Symp. Inform. Theory*, 2007.
- [31] ETSI. DVB, second generation framing structure, channel coding and modulation systems for Broadcasting, Interactive Services, News Gathering and other Broadband Satellite Applications Part 2: DVB-S2 Extensions (DVB-S2X);. European Standard, Oct. 2014. ETSI EN 302 307-2 V1.1.1.

- [32] ETSI. DVB, user guidelines for the second generation framing structure, channel coding and modulation systems for Broadcasting, Interactive Services, News Gathering and other Broadband Satellite Applications Part 2: DVB-S2 Extensions (DVB-S2X);. Technical Report, Nov. 2015. ETSI TR 102 376-2 V1.1.1.
- [33] Barry G. Evans, Paul T. Thompson, and Argyrios Kyrgiazos. Irregular beam sizes and non-uniform bandwidth allocation in HTS. In *AIAA Int. Commun. Satellite Systems Conf.*, pages 1–7, Florence, Italy, Sep. 2013.
- [34] S. Feizi, D. E. Lucani, and M. Médard. Tunable sparse network coding. *Int. Zurich Seminar on Communications*, February 2012.
- [35] H. Fenech, A. Tomatis, S. Amos, V. Soumholphakdy, and J. L. Serrano Merino. Eutelsat HTS systems. *International Journal of Satellite Communications and Networking*, 34(4):503–521, 2016.
- [36] H. Fenech, A. Tomatis, S. Amos, V. Soumholphakdy, and D. Serrano-Velarde. Future high throughput satellite systems. In *2012 IEEE First AESS European Conference on Satellite Telecommunications (ESTEL)*, pages 1–7, Oct 2012.
- [37] V. Firoiu, B. Adamson, V. Roca, C. Adjih, J. Bilbao, F. Fitzek, A. Masucci, and M. Montpetit. Network coding taxonomy. Technical report, Internet draft, draft-irtf-nwcr-network-coding-taxonomy-00.
- [38] S. Floyd, V. Jacobson, S. McCanne, C. g. Liu, and L. Zhang. A reliable multicast framework for light-weight sessions and application level framing. *ACM SIGCOMM Computer Communication Review*, 25(4):342–356, October 1995.
- [39] A. Ford, C. Raiciu, M. Handley, S. Barre, and J. Iyengar. Architectural guidelines for multipath tcp development. *IETF RFC*, 6182, 2011.
- [40] A. Ford, C. Raiciu, M. Handley, and O. Bonaventure. Tcp extensions for multipath operation with multiple addresses. *IETF RFC*, 6824, 2013.
- [41] Nikos Fotiou, Konstantinos V. Katsaros, George Xylomenos, and George C. Polyzos. H-pastry: An inter-domain topology aware overlay for the support of name-resolution services in the future internet. *Computer Communications*, 62:13 – 22, 2015.
- [42] Nikos Fotiou, Giannis F. Marias, and George C. Polyzos. Access control enforcement delegation for information-centric networking architectures. In *Proceedings of the Second Edition of the ICN Workshop on Information-centric Networking*, ICN '12, pages 85–90, New York, NY, USA, 2012. ACM.

- [43] J. Foust. Mega-constellations and mega-debris. <http://www.thespacereview.com/article/3078/1>.
- [44] C. Fragouli, J. Y. L. Boudec, and J. Widmer. Network coding: An instant primer. *ACM SIGCOMM Computer Communication Review*, 36(1):63–68, 2006.
- [45] G. Garramone, T. Ninacs, and S. Erl. On network coding for reliable multicast via satellite. In *Proc. of the 18-th Ka and Broadband Communications and 30-th AIAA International Communications Satellite Systems Conference*, Ottawa, Canada, September 2012. ICSSC.
- [46] G. Giambene, D. K. Luong, V. A. Le, and M. Muhammad. Network coding and mptcp in satellite networks. In *Proc. of the 2016 8th Advanced Satellite Multimedia Systems Conference and the 14th Signal Processing for Space Communications Workshop*, pages 5–7, 2016.
- [47] C. Hausl, O. Iscan, and F. Rossetto. Optimal time and rate allocation for a network-coded bidirectional two-hop communication. In Italy Lucca, editor, *Proc. of European Wireless*, April 2010.
- [48] T. Ho, M. Médard, R. Koetter, D. R. Karger, M. Effros, J. Shi, and B. Leong. A random linear network coding approach to multicast. *IEEE Transactions on Information Theory*, 52(10):4413–4430, October 2006.
- [49] V. Jacobson. Congestion avoidance and control. *SIGCOMM Comput. Commun. Rev.*, 18(4):314–329, August 1988.
- [50] Petri Jokela, András Zahemszky, Christian Esteve Rothenberg, Somaya Arianfar, and Pekka Nikander. Lipsin: Line speed publish/subscribe inter-networking. *SIGCOMM Comput. Commun. Rev.*, 39(4):195–206, August 2009.
- [51] M. Kaliski. Evaluation of the next steps in satellite high power amplifier technology: Flexible twtas and gan sspas. In *2009 IEEE International Vacuum Electronics Conference*, pages 211–212, April 2009.
- [52] Zoltán Katona, Federico Clazzer, Kevin Shortt, Simon Watts, Hans Peter Lexow, and Ratna Winduratna. Performance, cost analysis, and ground segment design of ultra high throughput multi-spot beam satellite networks applying different capacity enhancing techniques. *International Journal of Satellite Communications and Networking*, 34(4):547–573, Mar. 2016.
- [53] S. Katti, H. Rahul, H. Wenjun, D. Katabi, M. Médard, and J. Crowcroft. Xors in the air: Practical wireless network coding. *IEEE/ACM Transactions on Networking*, 16(3):497–510, June 2008.

- [54] F. P. Kelly and T. Voice. Stability of end-to-end algorithms for joint routing and rate control. *ACM SIGCOMM Computer Communication Review*, 35(2):5–12, April 2005.
- [55] R. Khalili, N. Gast, M. Popovic, U. Upadhyay, and J. y. Le Boudec. Non-pareto optimality of mptcp: Performance issues and a possible solution. *EPFL-REPORT-*, 177901, 2012.
- [56] R. Khalili, N. Gast, M. Popovic, and J. y. Le Boudec. Mptcp is not pareto-optimal: Performance issues and a possible solution. *IEEE/ACM Transactions on Networking*, 21(5):1651–1665, October 2013.
- [57] R. Khalili, N. Gast, M. Popovic, and J. y. Le Boudec. Opportunistic linked-increases congestion control algorithm for mptcp. *Internet draft, draft-khalili-mptcp-congestion-control-00*, February 2013.
- [58] A. Kyrgiazos, B. G. Evans, and P. Thompson. On the gateway diversity for high throughput broadband satellite systems. *IEEE Trans. Wireless Commun.*, 13(10):5411–5426, October 2014.
- [59] A. Kyrgiazos, P. Thompson, and B. G. Evans. System design for a broadband access terabits satellite for europe. In Italy Palermo, editor, *Proc. of the Ka band and Utilization Conference*, October 2011.
- [60] T. V. Lakshman and Upamanyu Madhow. The performance of tcp/ip for networks with high bandwidth-delay products and random loss. *IEEE/ACM Trans. Netw.*, 5(3):336–350, June 1997.
- [61] J. Lei and M. A. Vazquez-Castro. Joint power and carrier allocation for the multibeam satellite downlink with individual sinr constraints. In *IEEE Int. Conf. on Commun.*, pages 1–5, May 2010.
- [62] J. Lei and M. A. Vazquez-Castro. Multibeam satellite frequency/time duality study and capacity optimization. *J. of Commun. and Networks*, 13(5):472–480, Oct 2011.
- [63] S. Y. R. Li, R. W. Yeung, and N. Cai. Linear network coding. *IEEE Transactions on Information Theory*, 49(2):371–381, 2003.
- [64] E. Lier, D. Purdy, and K. Maalouf. Study of deployed and modular active phased-array multibeam satellite antenna. *IEEE Antennas and Propagation Magazine*, 45(5):34–45, Oct 2003.
- [65] Z. Liu, C. Wu, B. Li, and S. Zhao. Uusee: Largescale operational on-demand streaming with random network coding. In *Proc. of IEEE INFOCOM 2010*, pages 2070–2078.
- [66] J. Lizarraga, P. Angeletti, N. Alagha, and M. Aloisio. Flexibility performance in advanced ka-band multibeam satellites. In *IEEE Int. Vacuum Electronics Conf.*, pages 45–46, Apr. 2014.

- [67] M. Luby. Lt codes. In *Proc. of the IEEE Symposium on the Foundations of Computer Science*, pages 271–280, 2002.
- [68] M. Luby, A. Shokrollahi, M. Watson, and T. Stockhammer. Raptor forward error correction scheme for object delivery. Technical report, IETF RFC 5053, October 2007.
- [69] M. Luby, A. Shokrollahi, M. Watson, T. Stockhammer, and L. Minder. Raptorq forward error correction scheme for object delivery. Technical report, IETF RFC 6330, August 2011.
- [70] K. Maine, C. Devieux, and P. Swan. Overview of iridium satellite network. In *WESCON/'95. Conference record. 'Microelectronics Communications Technology Producing Quality Products Mobile and Portable Power Emerging Technologies'*, pages 483–, Nov 1995.
- [71] A. Mallet, A. Anakabe, J. Sombrin, R. Rodriguez, and F. Coromina. Multi-port amplifier operation for ka-band space telecommunication applications. In *2006 IEEE MTT-S International Microwave Symposium Digest*, pages 1518–1521, June 2006.
- [72] Saverio Mascolo, Claudio Casetti, Mario Gerla, M. Y. Sanadidi, and Ren Wang. Tcp westwood: Bandwidth estimation for enhanced transport over wireless links. In *Proceedings of the 7th Annual International Conference on Mobile Computing and Networking*, MobiCom '01, pages 287–297, New York, NY, USA, 2001. ACM.
- [73] T. Matsuda, T. Noguchi, and T. Takine. Survey of network coding and its applications. *IEICE Transactions*, 94-B(3):698–717, 2011.
- [74] R. Mauger and C. Rosenberg. Qos guarantees for multimedia services on a tdma-based satellite network. *Comm. Mag.*, 35(7):56–65, July 1997.
- [75] D. Mignolo, A. Ginesi, E. Re, A. Bolea Alamanac, P. Angeletti, and M. Harverson. Approaching Terabit/s satellite: a system analysis. In *17th Ka and Broadband Communications, Navigation and Earth Observation Conference*, pages 1–10, Palermo, Italy, Oct. 2011.
- [76] M.J. Miller and K.V. Buer. Flexible forward and return capacity allocation in a hub-spoke satellite communication system, December 25 2012. US Patent 8,340,016.
- [77] M.J. Miller and C.N. Pateros. Flexible capacity satellite communications system, September 25 2014. US Patent App. 13/666,112.
- [78] J. H. Møller, M. V. Pedersen, F. Fitzek, and T. Larsen. Network coding for mobile devices-systematic binary random rateless codes. In *Proc. of the International Conference on Communications*. ICC, 2009.

- [79] M. Mongelli, T. De Cola, M. Cello, M. Marchese, and F. Davoli. Feeder-link outage prediction algorithms for sdn-based high-throughput satellite systems. In *2016 IEEE International Conference on Communications (ICC)*, pages 1–6, May 2016.
- [80] FCC application Motorola Global Communications. Celestri multimedia leo system, June 1997.
- [81] M. Muhammad, G. Giambene, and T. de Cola. Channel prediction and network coding for smart gateway diversity in terabit satellite networks. In *Proc. of IEEE Globecom 2014 - Wireless Communications Symposium*, pages 8–12, Texas, December, 2014. Austin.
- [82] M. Muhammad, G. Giambene, and T. de Cola. Qos support in sgd-based high throughput satellite networks. *IEEE Transactions on Wireless Communications*, 15(12):8477–8491, Dec 2016.
- [83] K. Nichols, S. Blake, F. Baker, and D. Black. Definition of the Differentiated Services Field (DS Field) in the IPv4 and IPv6 Headers. RFC 2474 (Proposed Standard), December 1998. Updated by RFCs 3168, 3260.
- [84] OneWeb. Oneweb presentation, oct. 2016.
- [85] N. Porecki, G. Thomas, A. Warburton, N. Wheatley, and N. Metzger. Flexible payload technologies for optimising ka-band payloads to meet future business needs. In *19th Ka band Conference*, Sept. 2013.
- [86] R. Prior and A. Rodrigues. Systematic network coding for packet loss concealment in broadcast distribution. In *Proc. of the International Conference on Information Networking 2011*, pages 245–250, Barcelona, 26-28, January 2011. ICOIN2011).
- [87] T. Qi and Y. Wang. Energy-efficient power allocation over multibeam satellite downlinks with imperfect CSI. In *Int. Conf. on Wireless Commun. Signal Processing*, pages 1–5, Oct. 2015.
- [88] C. Raiciu, M. Handly, and D. Wischik. Coupled congestion control for multipath transport protocols. *IETF RFC*, 6356, 2011.
- [89] W. A. Sandrin. The butler matrix transponder. *COMSAT Technical Review*, 4(2):19–345, Fall 1974.
- [90] M. Schneider, C. Hartwanger, E. Sommer, and H. Wolf. The multiple spot beam antenna project "Medusa". In *2009 3rd European Conference on Antennas and Propagation*, pages 726–729, March 2009.
- [91] H. Shojania and B. Li. Parallelized progressive network coding with hardware acceleration. In *Proc. of the 15th IEEE International Workshop on Quality of Service*, pages 47–55, June 2007.

- [92] H. Skinnemoen. Gateway diversity in ka-band systems. In *Proc. of the 4th Ka-band Utilization Conference*, pages 2–4, Italy, November 1998. Venice.
- [93] P. Sourisse, D. Rouffet, and H. Sorre. Skybridge: a broadband access system using a constellation of leo satellites. <http://www.itu.int/newsarchive/press/WRC97/SkyBridge.html>.
- [94] R. Stewart. Stream control transmission protocol. *IETF RFC*, 4960, 2007.
- [95] M. A. Sturza and F. Ghazvinian. Teledesic satellite system overview. <http://tinyurl.com/gojlo8n>.
- [96] K. Sundararajan, D. Shah, M. Médard, S. Jakubczak, M. Mitzenmacher, and J. Barros. Network coding meets tcp: Theory and implementation. *Proceedings of the IEEE*, 99(3):490–512, March 2011.
- [97] G. Thomas, S. Laws, T. Waterfield, B. Gulloch, M. Trier, A. Montesano, and N. Gatti. Eutelsat quantum payload design and enabling technologies. In *3rd ESA Workshop on Advanced Flexible Telecom Payloads*, March 2016.
- [98] Y. Thomas, C. Tsilopoulos, G. Xylomenos, and G. C. Polyzos. Accelerating file downloads in publish subscribe internetworking with multi-source and multipath transfers. In *WTC 2014; World Telecommunications Congress 2014*, pages 1–6, June 2014.
- [99] Xenofon Vasilakos, Vasilios A. Siris, and George C. Polyzos. Addressing niche demand based on joint mobility prediction and content popularity caching. *Computer Networks*, 110:306 – 323, 2016.
- [100] O. Vidal, G. Verelst, J. Lacan, E. Albery, J. Radzik, and M. Bousquet. Next generation high throughput satellite system. In *2012 IEEE First AESS European Conference on Satellite Telecommunications (ESTEL)*, pages 1–7, Oct 2012.
- [101] F. Vieira, D. Lucani, and N. Alagha. Codes and balances: Multibeam satellite load balancing with coded packets. In *Proc. of the IEEE International Conference on Communications*, Ottawa, Canada, June 2012. ICC.
- [102] F. Vieira, S. Shintre, and J. Barros. How feasible is network coding in current satellite systems ? In *Proc. of the 5th Advanced Satellite Multimedia Systems Conference*, Cagliari, Italy, September 2010. ASMS and the *11th Signal Processing for Space Communications Workshop (SPSC)*.

- [103] M. Watson, T. Stockhammer, and M. Luby. Raptor fec schemes for fecframe. Technical report, Internet Draft, draft-ietf-fecframe-raptor-10, March 2012.
- [104] M. Werner. A dynamic routing concept for atm-based satellite personal communication networks. *IEEE Journal on Selected Areas in Communications*, 15(8):1636–1648, Oct 1997.
- [105] R. A. Wiedeman, A. B. Salmasi, and D. Rouffet. Globalstar: Mobile communications wherever you are. In *14th AIAA International Communications Satellite Systems Conference*, Sept. 1992.
- [106] D. Wischik, C. Raiciu, and M. Handley. *Balancing Resource Pooling and Equipoise in Multipath Transport*. ACM SIGCOMM, 2010.
- [107] G. Xylomenos, G. C. Polyzos, P. Mahonen, and M. Saaranen. Tcp performance issues over wireless links. *IEEE Communications Magazine*, 39(4):52–58, Apr 2001.
- [108] G. Xylomenos, X. Vasilakos, C. Tsilopoulos, V. A. Siris, and G. C. Polyzos. Caching and mobility support in a publish-subscribe internet architecture. *IEEE Communications Magazine*, 50(7):52–58, July 2012.
- [109] G. Xylomenos, C. N. Ververidis, V. A. Siris, N. Fotiou, C. Tsilopoulos, X. Vasilakos, K. V. Katsaros, and G. C. Polyzos. A survey of information-centric networking research. *IEEE Communications Surveys Tutorials*, 16(2):1024–1049, Second 2014.
- [110] D. Zhou and W. Song. *Multipath TCP for User Cooperation in Wireless Networks*. Springer, Brief in Computer Science, 2014.