



Characterizing Physical-Layer Transmission Errors in Cable Broadband Networks

Jiyao Hu, Zhenyu Zhou, and Xiaowei Yang, *Duke University*

<https://www.usenix.org/conference/nsdi22/presentation/hu>

This paper is included in the Proceedings of the
19th USENIX Symposium on Networked Systems
Design and Implementation.

April 4–6, 2022 • Renton, WA, USA

978-1-939133-27-4

Open access to the Proceedings of the
19th USENIX Symposium on Networked
Systems Design and Implementation
is sponsored by



جامعة الملك عبد الله
للعلوم والتقنية
King Abdullah University of
Science and Technology

Characterizing Physical-Layer Transmission Errors in Cable Broadband Networks

Jiyao Hu
Duke University

Zhenyu Zhou*
Duke University

Xiaowei Yang
Duke University

Abstract

Packet loss rate in a broadband network is an important quality of service metric. Previous work that characterizes broadband performance does not separate packet loss caused by physical layer transmission errors from that caused by congestion. In this work, we investigate the physical layer transmission errors using data provided by a regional cable ISP. The data were collected from 77K+ devices that spread across 394 hybrid-fiber-coaxial (HFC) network segments during a 16-month period. We present a number of findings that are relevant to network operations and network research. We estimate that physical-layer errors can contribute to 12% to 25% of packet loss in the cable ISPs measured by the FCC's Measuring Broadband America project. The average error loss rates of different HFC network segments vary by more than six orders of magnitude, from $O(10^{-6}\%)$ to $O(1\%)$. Users in persistently high-error-rate networks do not report more trouble tickets than other users.

1 Introduction

Reliable and high-speed Internet access is increasingly important to modern life, especially in a pandemic. According to [7], the number of broadband subscribers in the U.S. exceeded 105 million by the end of 2020. The availability and quality of broadband networks are of great policy concerns, as the U.S. government seeks to ensure an affordable and high-quality Internet service is provided to all [1].

In 2011, the Federal Communications Commission (FCC) launched the Measuring Broadband America (MBA) project to gain insight into the operational conditions of broadband networks [2]. The MBA project enlisted thousands of volunteers residing in ten U.S. ISPs and installed customized devices inside their homes. These devices send continuous measurement packets to estimate performance metrics such as packet loss rates, round trip latencies, and download/upload speeds of the volunteers' broadband networks. Similarly, much previous work measured and characterized different aspects of the last-mile broadband access

networks, including latency, loss, throughput, and availability [8, 17, 18, 25, 26, 28–30].

FCC's MBA project and previous work provide useful insight into how U.S. broadband networks perform. Among the metrics they gather, the packet loss rate is a particularly important Quality of Service (QoS) metric, as it affects TCP throughput as well as applications such as VoIP, live streaming, or multi-player online games. The communication quality of VoIP will significantly drop when the packet loss rate exceeds 1% [2]. In addition, the default TCP variant used by dominant operating systems, TCP Cubic [20], will reduce its sending rate after a packet loss.

However, the packet loss rates previous work measured have a severe limitation: they are end-to-end packet loss rates and do not separate the last-mile physical layer loss from other sources of packet loss. Packet loss comes from two main sources: error loss caused by the physical layer transmission errors and congestion loss caused by buffer contention at routers or switches. It is important to separate these two sources of packet loss for the following reasons.

First, physical layer packet loss is a direct indicator of how physical layer infrastructure functions, while other metrics, including latency, throughput, and end-to-end packet loss rates, are affected by multiple factors such as network capacity provisioning and router buffer management. Thus, physical layer loss can serve as a simple anomaly detector to network maintenance teams, while other metrics cannot.

Second, it is of great policy interest to monitor physical layer loss, as it is related to how well a broadband network is maintained. Broadband services in the U.S., while typically operated on existing telecommunications infrastructure (i.e., telephone or cable TV), are declassified from common carrier services [15]. Yet broadband Internet connections are increasingly becoming a public utility. Through continuous monitoring, policymakers can gauge how well the infrastructure is maintained without regulation and may consider appropriate policy adjustments if an unregulated broadband market leads to decreased quality of service.

Finally, understanding how much physical layer errors con-

*Zhenyu Zhou is now at Google.

tribute to end-to-end packet loss offers valuable insight into the design of congestion control algorithms and network simulations. A number of TCP variants, including Cubic [20], consider packet loss as a congestion signal. If physical layer error loss is common, we need to reexamine this assumption and possibly move away from such protocols to a more loss-agnostic one such as TCP BBR [14]. In addition, the design of a network protocol often uses simulations to evaluate the initial design. To conduct simulations, the designer often needs to configure a link's packet loss rate. To date, we do not have a clear understanding of how to configure the physical layer loss rate of a broadband link, but broadband networks are widely used in end-to-end connections. If we can separate the physical layer loss from the end-to-end packet loss, we can gain insight into how to build a physical layer error model and use it to conduct high-fidelity network simulations.

In this work, we aim to characterize packet loss caused by physical layer transmission errors. A regional cable ISP in the U.S. provides us physical layer performance data collected from 77K+ devices (primarily cable modems) every four hours from two disjoint geographical areas in a 16-month period.² The devices span across 394 hybrid-fiber-coaxial (HFC) network segments. Following operational practice, we refer to each HFC network segment as a fiber node (FN), as such a network segment terminates at a fiber optic node. The data we obtain include the number of unerrored, corrected, and uncorrectable DOCSIS [12] codewords a device sends since its last reboot. We develop techniques (§ 2) to use these codeword statistics as a proxy to understand the characteristics of the physical-layer transmission errors.

We make several observations that are relevant to network operations and research. First, we find that the average codeword error rate of an FN in our data spans six orders of magnitude, ranging from $1.3 \times 10^{-6}\%$ to 4.51%. The middle 80% of the FNs (excluding the top and bottom 10%) have average codeword error rates ranging from $9.53 \times 10^{-6}\%$ to $1.34 \times 10^{-3}\%$. We establish a relation between the codeword error rates in our data and the packet loss rates from FCC's MBA data, by assuming that the cable ISPs included in the MBA study have similar physical layer characteristics. We find that, for the five cable ISPs MBA monitors, even with a conservative estimate, 12% to 25% of the packet losses could come from the physical layer.

This finding has several ramifications. First, it establishes a baseline for a "normal" physical-layer error rate. If an ISP's packet loss rate significantly exceeds the baseline, it either indicates that there is an anomaly in the network infrastructure, or the congestion loss is high. Second, it challenges the assumption of loss-based congestion control protocols, as a significant percentage of packet loss can be attributed to physical layer errors even in wireline networks. Lastly, it suggests that comprehensive network measurements should use

packets of different sizes to measure packet loss, as codeword error loss is not negligible and packets of different sizes would be encoded in different numbers of codewords, resulting in different loss rates.

A second noteworthy finding is that we observe in a small number of FNs, all devices in those networks show codeword error rates exceeding 1% for months of time. Surprisingly, customers served by these devices do not make more trouble calls on average. In contrast, when customers who reside in the networks with a typical codeword error rate experience an error rate of the same value ($> 3\%$), they make nearly 15 times more daily customer calls. This discovery suggests that codeword error rates can reliably detect network faults in the absence of customer trouble tickets and ISPs should not solely rely on customer tickets to detect network maintenance issues. Based on this discovery, the ISP we collaborate with has developed an internal tool to periodically monitor codeword errors across its networks. In addition, the observation that codeword error rates of $> 1\%$ may persist for months suggests that congestion may not be the culprit when users experience poor application performance.

Finally, we analyze how codeword error rates change before and after COVID-19 and find that the error rates are not impacted by the increase in traffic loads. We find that the codeword error rates of devices in the same FN are more correlated when their codeword error rates are high. We also study how weather impacts codeword error rates. The results show that extremely high ($> 95^\circ F$) or low ($< 15^\circ F$) temperatures increase codeword error rates, while the types of precipitation (e.g., freezing rain, snow) tend to cause outages than increase codeword error rates.

A limitation of this work is that our findings are based on data from one ISP. That being said, the ISP that provides us the data follows standard industry practices and uses standard Cable Modem Termination System (CMTS) equipment from dominant vendors. While different ISPs may choose CMTS modulation profiles to overcome specific radio frequency (RF) impairments at the cost of potentially reduced capacity, we are not aware of other reasons that will cause the overall physical layer loss characteristics of one ISP to differ from those of others. We release the code and part of the data used for this study.³

To the best of our knowledge, this work is the first large-scale and public study on the characteristics of physical-layer transmission errors of cable broadband networks. We make three main contributions. First, we characterize the physical-layer transmission errors of 394 HFC network segments and establish the relationship between physical-layer transmission errors and packet loss measured by FCC's MBA project. Second, we show that physical-layer transmission errors can indicate network faults in the absence of trouble tickets. Finally, we show that codeword errors are not impacted by

²Due to our non-disclosure agreement, we cannot disclose the locations of the devices or the name of the ISP.

³<https://github.com/zhenyu-zhou/pnm-loss-nsdi22>

traffic loads or types of precipitation, but tend to increase in extremely cold or hot weather.

2 Methodology

In this section, we describe how we estimate the physical layer transmission errors and how we relate them to upper layer packet loss.

2.1 DOCSIS Codeword

The data items used in this study are the DOCSIS codeword statistics. Before we describe the data, we first describe what DOCSIS codewords are and how they impact upper layer packet loss. A codeword is a cable modem's basic transmission unit at the physical layer. The cable modems used in this study are DOCSIS 3.0 modems. DOCSIS 3.0 uses Forward Error Correction (FEC) to detect and correct errors at the physical layer. A codeword includes a data section and an FEC parity check section. In DOCSIS 3.0, each codeword is generated using a Reed Solomon (RS) encoder. The size of a codeword can vary from 18 bytes to 255 bytes, containing k data bytes and $2T$ parity check bytes. An RS codeword with $2T$ parity check bytes can correct up to T byte or $8T$ bit errors [13]. Both the data length k and the parity check length $2T$ of a codeword are vendor and configuration dependent. Typically, a CMTS vendor specifies a default setting of k and T for a long codeword and a short codeword. Cable operators can choose different settings, but the current industry practice is to use the default settings chosen by vendors.

Most of our data are collected from CMTS devices manufactured by a dominant vendor in the U.S. As an example, the default setting for our data includes two codeword lengths: one long codeword and one short codeword. The long codeword has a data length k of 200 bytes and a parity check byte length of $2T = 30$ bytes. The long codeword is able to correct 15 bytes of errors. Similarly, the short codeword has a data length $k = 99$ bytes, and a parity check byte length $2T = 10$ bytes. It is able to correct 5 bytes of errors in a codeword.

When a cable modem receives a data frame from an upper layer protocol such as Ethernet, if the data frame fits into a long or a short codeword, it will transmit the data frame using one codeword. Otherwise, the modem will use multiple codewords to transmit the data frame. A cable modem will at most use one short codeword at the end of a data frame to transmit it. If a data frame does not fit exactly into multiple codewords, the cable modem will use padding bytes at the last codeword. Figure 1 shows an example of how a cable modem encodes an Ethernet MAC frame into multiple codewords.

As bit errors at the physical layer tend to be bursty, a cable modem uses a scrambler to permute the content of a codeword before transmission, following a pre-defined pseudo-random pattern. The receiving end, the CMTS, will reverse the permutation before decoding the received data. Therefore, bursty errors in transmitted signals become random errors in the unscrambled codewords.

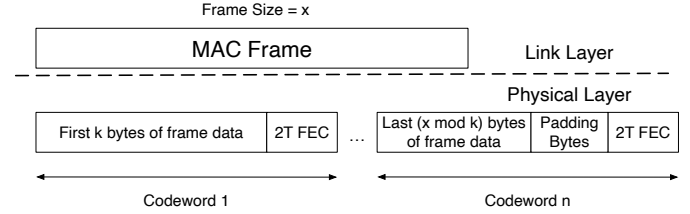


Figure 1: A link layer MAC frame is encoded by multiple codewords at the physical layer. Each codeword has a data section and an FEC section.

2.2 Codeword Error Rate

We refer to the ISP that provides us data as AnonISP. We now describe how we compute the codeword error rate using data collected by AnonISP. AnonISP collects the data through their Proactive Network Management (PNM) platform [11], which is part of DOCSIS's design. It aims to help cable ISPs troubleshoot and diagnose their networks. With PNM, an ISP can collect various performance metrics from both cable modems and a CMTS, including codeword counters, signal transmission power (TX power), and signal to noise ratio (SNR).

For each cable modem, AnonISP collects the total number of unerrored codewords sent by a cable modem, the number of uncorrectable codewords that fail FEC, and the number of codewords corrected by FEC periodically. All numbers are cumulative since the modem's last reboot. From DOCSIS 3.0's specification [13], uncorrectable codewords are discarded without link-layer retransmission. For correctable codeword errors, they do not manifest them as upper-layer packet discards. Therefore, we focus on the uncorrectable codeword errors. Without specific clarification, in this work, we use *codeword errors* to refer to uncorrectable codeword errors.

We estimate the average codeword error rate $P(e)$, where e denotes packet loss events, of a cable modem as the number of *uncorrectable* codewords divided by the *total* number of codewords a CMTS receives in each collection period: *uncorrectable/total*. Note that in our data, codewords are of two different lengths. The codeword error rate of a long codeword and a short codeword could be different, but we can still estimate the average codeword error rate without knowing the distribution of short or long codewords. Formally, let $P(e|l)$ denote the probability of a long codeword error rate and $P(e|s)$ be the probability of a short codeword error rate. Let $P(l)$ and $P(s)$ be the probability of a long or short codeword occurring in the data stream, respectively. We can compute the average codeword error rate $P(e)$ as follows:

$$\begin{aligned}
 P(e) &= P(e|l)P(l) + P(e|s)P(s) \\
 &= \frac{\text{uncorrectable long}}{\text{long}} P(l) + \frac{\text{uncorrectable short}}{\text{short}} P(s) \\
 &= \frac{\text{uncorrectable long}}{\text{long}} \times \frac{\text{long}}{\text{total}} + \frac{\text{uncorrectable short}}{\text{short}} \times \frac{\text{short}}{\text{total}} \\
 &= \frac{\text{uncorrectable}}{\text{total}}
 \end{aligned}$$

2.3 Codeword Errors vs. Packet Loss Rates

We aim to understand how physical layer transmission errors affect end-to-end packet loss. We ask this question: *how much do physical layer transmission errors contribute to higher-layer packet loss?* Since FCC’s MBA project measures packet loss on broadband networks, if we can establish the relationship between codeword errors at the physical layer and packet loss measured by the MBA project, then we can estimate how much packet loss is caused by physical layer errors. To do so, we make the assumption that the physical layer loss characteristics of AnonISP’s networks are representative of those of U.S. cable broadband networks. With this assumption, we can correlate the network-layer packet loss rates from the MBA project with codeword error rates observed in our data.

There are three challenges in establishing the correlation. First, a packet has a variable length so that it may be encoded in multiple codewords. Hence, there does not exist a one-to-one correspondence between the codeword error rate and the packet loss rate. Fortunately, FCC’s MBA project uses short UDP ping packets with packet length set to 62 bytes [4] to continuously monitor the packet loss rates. Such a packet will be encoded using one short codeword under common CMTS configurations. Therefore, if we assume cable broadband networks operate in similar physical conditions, then the short codeword error rate will correspond to the packet loss caused by physical layer errors in the MBA project.

Second, the codeword error rate we measure is the average codeword discard rate that includes both short and long codewords, while the FCC MBA project uses only short UDP packets that correspond to short codewords for measuring packet loss. To address this challenge, we analyze whether the average codeword error rate is an over- or under-estimate of the short codeword error rate. According to the common CMTS configurations, for a long codeword to become uncorrectable, more than 120 bits out of 200 bytes must be errored. For a short codeword to become uncorrectable, more than 40 bits out of 99 bytes must be corrupted. Since a long codeword is roughly twice the size of a short codeword, and the number of FEC bits in a long codeword is three times that in a short codeword, the long codeword should have a much lower error rate than the short one, assuming the bit error rate in a long or a short codeword is the same. Therefore, the average codeword error rate in our data is a lower bound to the short codeword error rate. In other words, if the cable networks operate in similar physical conditions, the average codeword error rate we measure is a lower bound to the physical layer error rate from FCC’s MBA project, since the measurement project only uses short UDP packets.

Third, the codeword statistics we obtain only include the upstream channels. That is, we only observe the codeword errors from a customer’s device to the ISP’s cable headend. However, packet loss measured by FCC’s MBA project is bi-directional. To reconcile the difference, we use the up-

stream transmission errors as a lower bound to bi-directional transmission errors and an upper bound to downstream transmission errors. DOCSIS’s downstream channels operate at higher frequencies than upstream channels [13], while RF noises concentrate on the lower RF range. In § 4.1, we show how the codeword error rate decreases as a channel’s frequency increases. Therefore, downstream channels should have lower codeword error rates than upstream channels and we can use the upstream transmission errors to upper-bound downstream transmission errors.

3 Datasets

Next, we describe our datasets and data cleansing steps.

AnonISP Data At the time the data were collected, AnonISP uses three upstream channels and sixteen downstream channels as data channels in their networks. Each channel is of 6MHz width. Our data include the upstream codeword statistics only. At each data collection time point, AnonISP collects several metrics for each upstream channel, including SNR, cumulative values of the number of unerrored codewords, the number of corrected codewords, and the number of uncorrectable codewords each cable modem sends to a CMTS since it reboots, and the signal transmission (TX) power of a cable modem. The data is collected from 01/06/2019 to 03/03/2020 and from 03/24/2020 to 04/17/2020. The data are collected every 4 hours. In total, the data come from 77,696 devices and span 394 fiber optical nodes. On average, each fiber node has 197 devices. In total, we have collected $\sim 139\text{M}$ data points for each upstream channel. We call this dataset the codeword dataset.

In the codeword dataset, each data point contains the data from three upstream channels. If these three channels send fewer than 2,000 codewords in total between the current and its previous data collection point, which means the three channels send less than 200KB of data during the last 4 hours, we will consider the current data point invalid as the loss statistics may be distorted because of too few numbers of codewords.

In addition, it is possible that at a data collection point, we fail to retrieve data from a device. Since our data are collected every 4 hours, if we observe that the time interval between two adjacent data points is $4 \times (1 + x)$ hours (where x rounds to an integer), we will insert x empty placeholder data points in the data stream. These empty placeholders indicate that we fail to retrieve data at those time points. We infer empty placeholder data points and refer to them as *missing data*. If all three channels’ data are missing, we will count this data point as a *missing* data point, which can indicate that the network is unavailable. If only one or two channels have missing data, we will discard this data point, because we often combine the three channel’s data for our analysis.

Among all of the data points ($\sim 139\text{M}$), we discard $\sim 33\text{M}$ data points and obtain $\sim 106\text{M}$ valid data points. In addition, we infer $\sim 11\text{M}$ missing data points when a collection point fails to collect any data. Among the 33M discarded data

points, $\sim 9\text{M}$ of the data points are discarded due to missing partial channel data, while $\sim 24\text{M}$ of the discarded data points have fewer than 2,000 codewords. We use the valid data points and the missing data points for our analysis in this paper.

Besides the codeword dataset, AnonISP also provides us the customer call trouble tickets from the same group of devices during the same time periods. Each trouble ticket includes the call time, the description of the issue that triggered the customer call, and how AnonISP resolved the issue.

FCC Data from MBA Project To understand the relation between codeword error rates and packet loss rates, we compare our data with the FCC data obtained from the MBA project [2]. The FCC data are continuously collected from thousands of users all over the United States since Jan 2011 and are available to the public. FCC deployed whitebox measurement devices in volunteers' homes. The volunteers are distributed across 10 wireline broadband providers. The measurement devices continuously send UDP packets to target test nodes to measure packet loss rates. If a device does not receive a response packet within three seconds, it labels the packet as lost. The devices follow the Poisson distribution to send probe packets over a fixed interval of one hour [4]. We use the FCC data collected from the same period as our data. The FCC data contain several broadband technologies, including DSL, Cable, and FTTH. Since our data are from cable networks only, we only analyze the data collected from cable networks in the FCC data, and leave a comparison among different physical layer techniques as future work.

Weather Data We collect weather data that overlaps with the codeword dataset in time and location to study how weather affects physical layer transmission errors. We use the IBM Weather Data APIs [6] to collect the hourly weather conditions given a time period and the zip code each device in our dataset belongs to. Each weather data record includes the basic weather metrics such as temperature, atmospheric pressure, and humidity. It also contains the description of the current weather type, such as Light Rain or Snow.

Ethical Considerations Prior to obtaining data from AnonISP, we consulted with our organization's IRB and obtained their permission to conduct this research. The MAC address and account number provided by AnonISP are hashed values. All the statistics in our data are performance monitoring metrics generated by the devices. For each customer trouble ticket, it only records the time of the trouble call, the hashed account number to match the monitoring metrics, the description of the issue, and the action of the ISP. So there is no personal identification information included in our data.

4 Physical Layer Loss vs. Overall Loss

In this section, we study how the physical layer errors look like using the codeword dataset. We compare the physical layer errors with packet loss observed in the FCC data, aiming

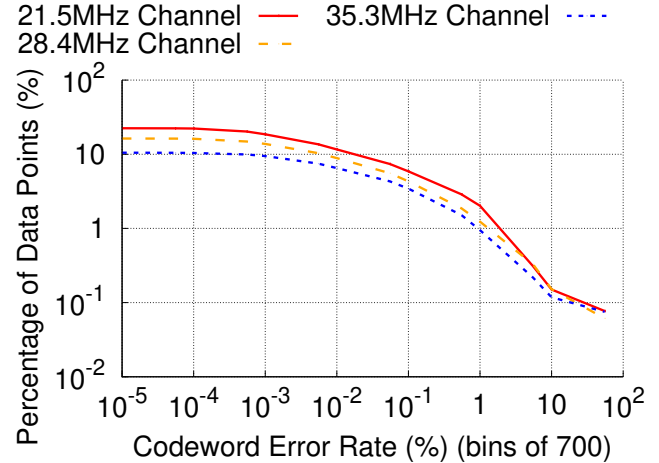


Figure 2: **The complementary cumulative distribution of the codeword error rate in each upstream channel. The error rate decreases when the channel frequency increases.**

to answer the question: *What is the relationship between the physical layer codeword error rate and the end-to-end packet loss rate?*

4.1 Codeword Errors in Different Channels

Our codeword data are collected from three upstream channels in AnonISP's HFC networks. The three upstream channels send RF signals with 21.5MHz, 28.4MHz, and 35.3MHz center frequency, respectively. The downstream channels will each use a higher center frequency, ranging from 54 MHz to as high as 1000 MHz. Figure 2 plots the complementary cumulative distributions of all three channels' codeword error rates, respectively. Each data point is computed as the number of uncorrectable codewords divided by the total number of codewords a device sends to a CMTS over the 4-hour data collection period. Both the x-axis and y-axis are in log-scale.

From Figure 2, we can see that the majority of the data points have no or few codeword errors, as seen in the flat sections at the beginning of the lines. At least 75% of the data points in each channel have no uncorrectable codewords. However, for all three channels, the curves start to drop after the error rate exceeds 10^{-4} , suggesting that a small fraction of lossy periods contribute to the majority of codeword errors. In particular, more than 1% of the data points have codeword error rates exceeding 1%; and more than 0.1% of the data points have codeword error rates exceeding 10%. The average codeword error rates in the three channels are 0.11%, 0.08%, and 0.06%, respectively. Furthermore, the channels with a higher center frequency have lower error rates, consistent with the operational knowledge that lower frequency channels are more prone to RF interference.

In the DOCSIS design, a cable modem will switch to a different upstream channel when one upstream channel does not work expectedly. We study how codeword errors in the upstream channels are correlated. That is, for each data point, we investigate whether the three channels show similar codeword

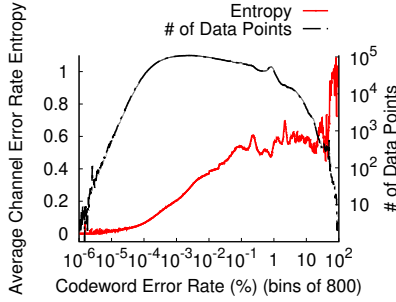


Figure 3: The relationship between the average channel error rate entropy and the codeword error rate together with the number of data points in each bin.

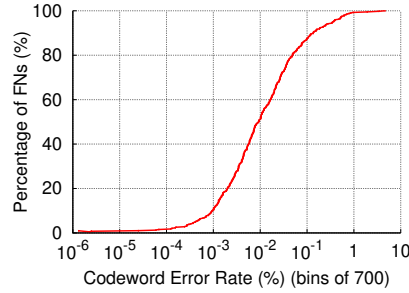


Figure 4: CDF of the average codeword error rate of an FN.

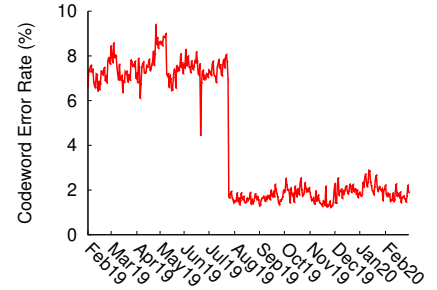


Figure 5: The daily-average codeword error rate of the FN with the highest average codeword error rate in our data.

error rates or the codeword error rates of the three channels differ by a lot. To quantify the similarity of codeword errors in each data point, we compute the channel error entropy S , where S is defined as the error rate entropy of each data point. That is, let s_i be the uncorrectable codewords sent via channel i divided by the total number of uncorrectable codewords across all three channels in each data point. We compute $-\sum_{i=1}^3 s_i \ln s_i$ for each data point. A higher channel error entropy value indicates a modem has a lower error rate variation among the three channels for this data point. If all the uncorrectable codewords are from one single channel, then S will be 0. In contrast, if the three channels have the same codeword errors, the value of S will achieve its maximum.

Figure 3 shows the relationship between the average channel error rate entropy and the codeword error rate. We divide the codeword error rate values into 800 bins and calculate the average error rate entropy of data points in each bin. This figure shows that the average channel error rate entropy increases as the codeword error rate increases, suggesting that when codeword errors in one upstream channel occur, they are likely to occur in other upstream channels as well. So DOCSIS's upstream channel switching algorithm may be insufficient to avoid codeword errors.

Figure 3 shows that codeword errors are highly correlated in three upstream channels when the codeword error rate exceeds 0.1%. Therefore, without specific mentioning, we will use the number of combined codewords and the number of combined codeword errors from three channels in each data point for our analysis in the rest of this paper. The average codeword error rate from the combined channels is 0.088%, while 98.68% data points have a codeword error rate $< 1\%$.

Takeaways: Codeword errors occur infrequently, and a small percentage of lossy periods contribute to most of the codeword errors. Higher frequency channels have lower codeword error rates, and when codeword errors occur, they tend to occur in all upstream channels. We show more examples of raw codeword error rates in Appendix A.

Next, we investigate *whether the devices in different fiber*

optic nodes (FNs) will have different codeword error rates. To do so, we compute the codeword error rate of each device over the 16-month data collection period. We then compute the average codeword error rate of each FN by averaging the codeword error rates of all devices in the FN.

Figure 4 shows the CDF of the average codeword error rate among 394 FNs in our data. The x-axis is again in log-scale. We observe that the average codeword error rate differs significantly among different FNs. In our data, there are three FNs that have an average codeword error rate higher than 1%, and 46 FNs have an average codeword error rate between 0.1% to 1%. We define the FNs with $> 1\%$ error rates as unhealthy FNs, the FNs with 0.1% - 1% error rates as alarming FNs, and the remaining 345 FNs as healthy FNs. The thresholds 1% and 0.1% are chosen according to operational experience suggested by AnonISP. The healthy FNs constitute 87.6% of the FNs seen in our data and the codeword error rate averaged over those FNs is 0.0179%. They contribute to only 18.79% of the total codeword errors. In contrast, the alarming FNs are 11.68% of the FNs seen in our data and their average codeword error rate is 0.352% and they contribute to 42.67% of the total codeword errors. The unhealthy FNs are 0.761% of the total FNs. Their average codeword error rate is 3.778% and they contribute to 38.54% of the total codeword errors.

We are interested in understanding why certain FNs have such high codeword error rates. Figure 5 shows the daily codeword error rate from the FN with the highest average codeword error rate in our data. The average error rate is over 7% before July 2019, and then it decreases to 2%, suggesting that AnonISP repaired some problems in this node. However, even after this repair event, this FN still has a nearly 2% daily codeword error rate, and it exists till the end of our data. We are informed by AnonISP that this FN is affected by issues in the network hardware. Its maintenance team is aware of this persistent problem, but it either cannot repair the issue for some reason (e.g., waiting on permits, access and/or restrictions, etc.) or has deprioritized the repair for some reason (e.g., it is a known but uncorrectable cause).

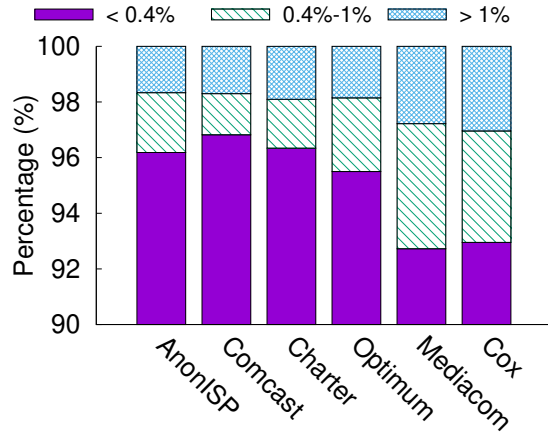


Figure 6: This figure shows the percentage of data points whose packet loss rate was less than 0.4%, between 0.4% to 1%, and greater than 1% for each cable ISP in the FCC data, together with AnonISP’s loss rate we measure.

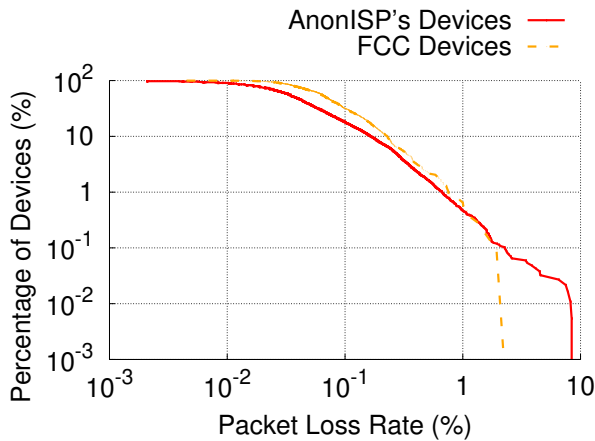


Figure 7: This figure shows the complementary cumulative distribution of the average packet loss rate of each device in the FCC data and in the AnonISP’s data we measure, respectively.

Takeaways: Codeword error rates in different HFC network segments vary significantly. Some network segments may experience persistent high codeword error rates ($> 1\%$). There are 87.6% healthy FNs in our data and they contribute to 18.79% of the total codeword errors. The 12.4% alarming and unhealthy FNs contribute to 81.21% of codeword errors seen in our data.

4.2 Comparison to FCC data

A key question this work aims to answer is how much physical layer error loss contributes to end-to-end packet loss. We compare our data with the FCC data collected by the MBA project to gain insight into this question. The FCC data measures the packet loss rates in different types of networks, including Cable, FTTH, and DSL. We only used the data collected from cable ISPs in the FCC dataset.

The FCC dataset does not include data from AnonISP, which prevents us from comparing the physical layer loss of

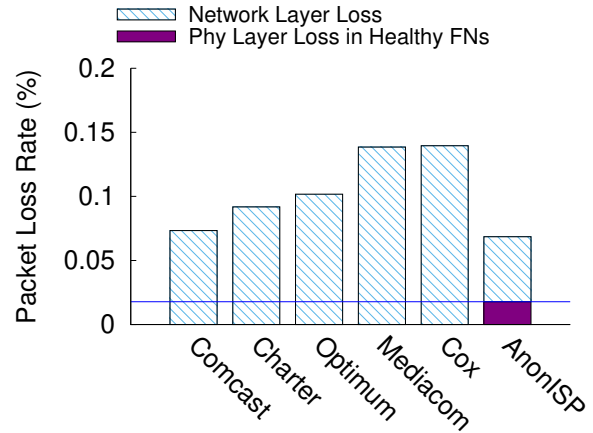


Figure 8: The average packet loss rate in each cable ISP, together with the packet loss rate we measure in AnonISP and the physical layer error rate among the health FNs in AnonISP.

AnonISP’s networks directly with end-to-end packet loss. To address this challenge, we design an experiment to approximate the FCC’s packet loss measurement for AnonISP. We deploy a measurement node on a vantage point that is close in router hops to AnonISP’s networks. The vantage point sends ICMP echo request packets to all 18,772 pingable modems located in AnonISP’s cable networks in the data collection regions periodically. For each modem, we send ~ 250 ICMP packets per hour. We run this experiment from 11/03/2021 to 11/11/2021. The measurement node sends 914M packets in total. The FCC dataset has been cleansed to exclude data points with high loss rates ($> 10\%$) and high RTTs [5]. We applied to our measurement results the same data cleansing script as applied to the FCC data.

Figure 6 shows the percentages of end-to-end packet loss rates in different ranges for the FCC dataset, together with the end-to-end packet loss rate we measure for AnonISP. We note that the ICMP packet loss rate observed in AnonISP is comparable to the packet loss rates observed in the FCC measurement. Similar to our measurement results, most data points in the FCC data suffer no or few packet losses and the majority of packet loss comes from a small percentage of lossy periods. For example, around 97.03% data points in our data have less than 0.4% loss rate, and 1.73% data points have loss rates between $[0.4\%, 1\%]$, while in Comcast’s data, 96.43% data points have less than 0.4% packet loss rate, and 1.72% data points have packet loss rates between $[0.4\%, 1\%]$.

We compute the average packet loss rate for each device in the FCC data and that for each device in our measurement. There are a total of 1,073 devices installed in five cable ISPs in the FCC measurement. In contrast, we measure 18K+ devices. Figure 7 shows the results. Since we measure more devices, we see a wider range of packet loss rates in our measurement than that in the FCC data.

Figure 8 shows the average packet loss rate in each cable ISP, together with the average physical layer error rate

among the healthy FNs seen in our data. We compare the FCC packet loss rate with the physical layer codeword error rate from healthy FNs only because FCC’s volunteers are sparsely located. The alarming and unhealthy FNs account for only 12.44% of the 394 FNs in our data. Therefore, there may or may not be any devices located in those outlier FNs in the FCC study. As we aim for a lower-bound estimate regarding the physical layer’s contribution to end-to-end packet loss, we exclude the FNs with the high error rates in our data from the comparison. We assume that the physical layer error rates in those ISPs’ healthy FNs are the same as those in our data (0.0179%) since we think our data is representative of the nature of cable networks. Based on this assumption, we see that at least from 12% to 25% packet loss seen in the FCC data could have come from physical layer errors.

Meanwhile, we estimate how much ICMP packet loss from our own measurement could come from physical layer transmission errors. If we assume that our sampled devices do not include any devices in the alarming or unhealthy FNs, then 26.1% of the packet loss seen in AnonISP can be attributed to physical layer codeword errors. However, this estimate may be overly conservative, as we receive ICMP echo replies from more than 20% of all devices in our codeword dataset. If we assume that we have representatively sampled devices from both healthy and alarming FNs, and since the average codeword error rate among the healthy and alarming FNs in our data is 0.0551%, and the average ICMP packet loss rate we measure is 0.0686%, then 80.3% of packet loss in our measurement could have come from physical layer transmission errors. However, since we cannot establish a one-to-one correspondence between a device we ping and a device we see in the codeword dataset, due to the anonymization procedure applied to the data, we cannot conclude whether the devices we ping are a representative subset of devices from the healthy and alarming FNs. Therefore, we prefer to use 26.1% as a safer lower bound.

Our estimate presents how much physical layer errors contribute to end-to-end packet loss. This is a lower bound estimate for the following reasons. First, the baseline codeword error rate we compute is the average between long and short codewords, while the FCC measurement packets only use short codewords (§ 2.3). The short codeword’s error rate is higher than the average codeword error rate due to the encoding scheme. Second, the baseline codeword error rate only includes codeword errors in the upstream channels, while the packet loss measurement is affected by both upstream and downstream errors. Third, we have conservatively excluded all alarming and unhealthy FNs in our codeword error rate calculation, while the FCC measurement may include such nodes. Finally, the baseline codeword error rate only includes the network segment from a cable modem to a CMTS, while the physical layer errors in the entire end-to-end paths can contribute to packet loss in the FCC measurement.

We also note that different cable ISPs have very different

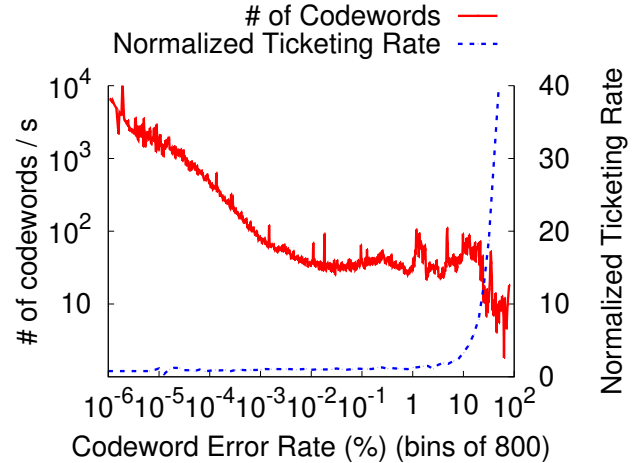


Figure 9: This figure shows how codeword error rate affects the number of codeword sent per second by a device and the normalized customer ticketing rate. The number of data points in each bin is the same as shown in Figure 3.

packet loss rates. Among the five ISPs in the FCC data, Comcast shows the lowest average packet loss rate: 0.073% with a standard deviation of 0.486%. In contrast, the average packet loss rates of Mediacom and Cox are 0.138% (with a standard deviation of 0.621%) and 0.140% (with a standard deviation of 0.619%), respectively. They are almost two times higher than Comcast’s average packet loss rate. AnonISP shows the lowest packet loss rate among the six ISPs. Its packet loss rate is slightly lower than Comcast’s. We speculate that this is because we place the measurement node close to the cable modems and our measurement packets have shorter RTTs than FCC’s measurement packets. Therefore, they encounter fewer congestion and physical layer transmission error events.

Takeaways: We show that 12% to 25% of the packet loss measured by FCC’s MBA project on cable ISPs could have come from physical layer errors. This result suggests that physical layer errors in cable networks play a non-negligible role in end-to-end QoS. Network research and operations should take this source of packet loss into account.

5 Analysis of User Behavior

In this section, we investigate how codeword errors affect user behavior. We use the amount of data sent by customer devices and the customer reported trouble tickets to quantify user behavior and study how they change when codeword error rates change.

5.1 Impact on Usage

We use the number of codewords in each data point to estimate the amount of data users sent, because application data will be sent using codewords. We divide the range of codeword error rates into 800 bins. For each bin, we calculate the total number of codewords from the data points falling into this bin and normalize it by the period of time covered by the data points in the bin. With this computation, we obtain the

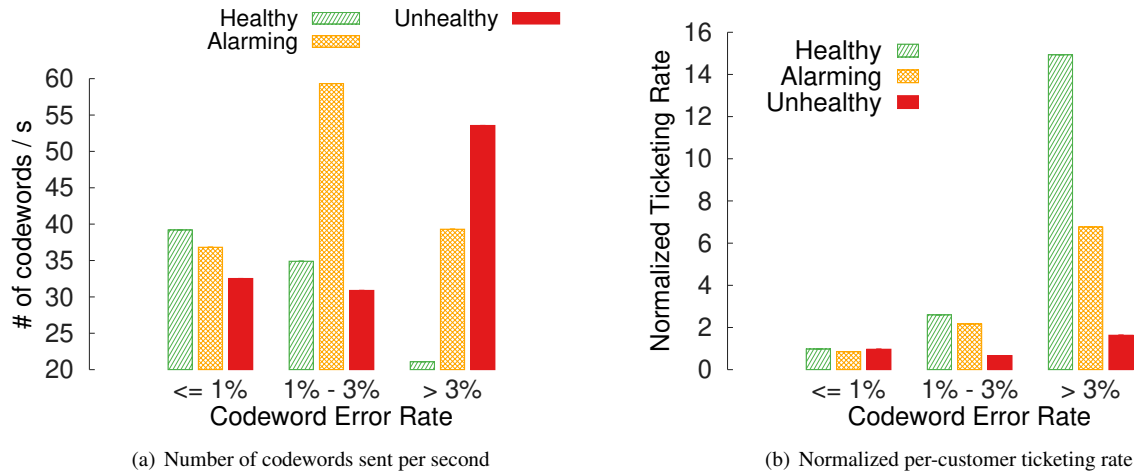


Figure 10: This figure shows how users or devices in different types of FNs behave when experiencing different codeword error rates.

number of codewords sent by a device per second when a specific codeword error rate occurs.

Figure 9 shows the relationship between the codeword error rate and the number of codewords sent per second by a customer device. We can see that the number of codewords sent per second by a user device decreases as the codeword error rate increases from $10^{-6}\%$ to $10^{-2}\%$. It plateaus between $10^{-2}\%$ and 1%, and sharply decreases when the codeword error rate increases beyond 10%. The data between 1% and 10% error rates are jagged. One plausible explanation is that loss rates within this range will significantly impact user experience [2, 3], and users or applications may react to the adverse conditions by multiple retries, leading to a fluctuated data rate.

5.2 Impact on Customer Trouble Tickets

When a customer has poor QoE, she may call her ISP's customer service to report the issue. Therefore, customer tickets are a good indicator of network problems and also reflect customer experience [22]. We study how customer reported trouble tickets are affected by codeword errors. We define the *ticketing rate* as the average number of trouble tickets each customer reports in a unit time. We compute a baseline ticketing rate by computing the average number of tickets reported by each customer in a unit time. We define a *normalized ticketing rate* as a ticketing rate divided by the baseline ticketing rate.

Figure 9 shows the relationship between the normalized ticketing rate and the codeword error rate. Similarly, we divide the range of codeword error rates into different bins and compute the normalized ticketing rates for data points that fall into each bin. The customer ticketing rates remain stable until the codeword error rate exceeds 1%. It increases fastly after that. In extreme cases, when the codeword error rate exceeds 50%, the customer ticketing rates increase by more than 40

times compared to the baseline ticketing rate.

5.3 Conditioned User Behavior

Next, we investigate how codeword error rates impact a user's behavior for users located in different network environments. From our study in § 4.1, we show that different FNs can have drastically different codeword error rates. We classify the FNs in our data into three types based on their average codeword error rates: healthy ($< 0.1\%$), alarming ($[0.1\%, 1\%]$), and unhealthy ($> 1\%$). We divide the codeword error rates into three ranges $< 1\%$, $[1\%, 3\%]$, and $> 3\%$ and examine how user behavior in different types of FNs varies in different codeword error ranges. Specifically, we compute the number of codewords sent per device and the normalized ticketing rate for each codeword error range for healthy, alarming, and unhealthy FNs, respectively.

Figure 10(a) and 10(b) show the results. For users in healthy FNs, their data usage decreases and their normalized ticketing rate increases as the codeword error rate increases. This trend is consistent with the general trends shown in Figure 9.

The usage and ticketing rate patterns in the alarming and unhealthy FNs are somewhat counter-intuitive. Customers in unhealthy FNs report much fewer tickets on average, when their networks show a $> 3\%$ loss rate. In contrast, for the customers in healthy FNs, when the codeword error rate is larger than 3%, the probability of a customer reporting a ticket will increase by 14.93 times. Instead, the customers in unhealthy FNs increase their data usage when the error rate exceeds $> 3\%$, suggesting that they or their applications attempt to use retransmissions or redundant transmissions to overcome packet loss. For customers in alarming FNs, their behavior is even more puzzling. They increase their data usage when the error rate is in the $[1\%, 3\%]$ range and decrease the usage when it exceeds 3%. One plausible explanation is that the customers in those FNs would attempt retry or retransmission

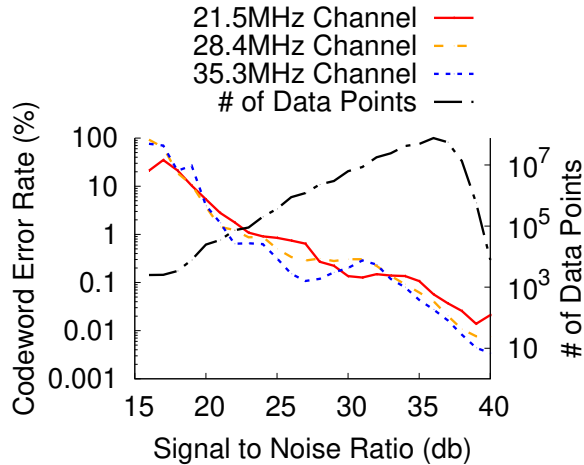


Figure 11: The correlation between the codeword error rate and the SNR together with the number of data points with respect to an SNR value.

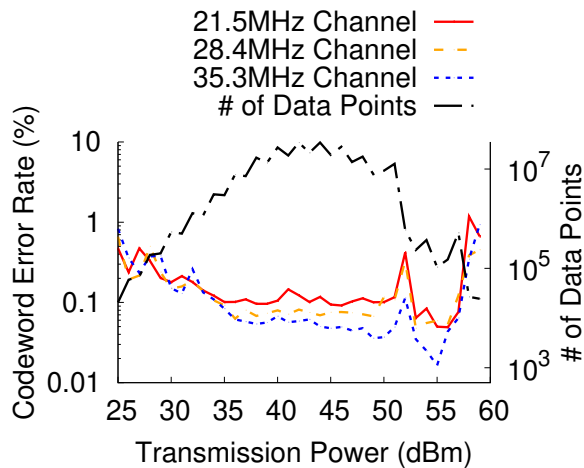


Figure 12: The correlation between the codeword error rate and the TX power together with the number of data points with respect to a TX Power value.

first when the network conditions slightly worsen, but will give up using the networks when the network conditions are significantly worse than what they are used to.

Takeaways: Users generally report more trouble tickets when the codeword error rate increases. However, users belonging to an FN with a consistently high codeword error rate have a higher tolerance for packet loss. This result indicates that network operators should continuously monitor the codeword error rates of their networks. Lack of trouble tickets alone is not a reliable indicator of good network conditions.

6 What Affects Codeword Error Rate?

In this section, we study what factors impact codeword errors. We examine how other PNM metrics (SNR and TX power) correlate with codeword error rate, how the traffic load increases after COVID-19 and different weather conditions affect the codeword errors in an HFC network, and how

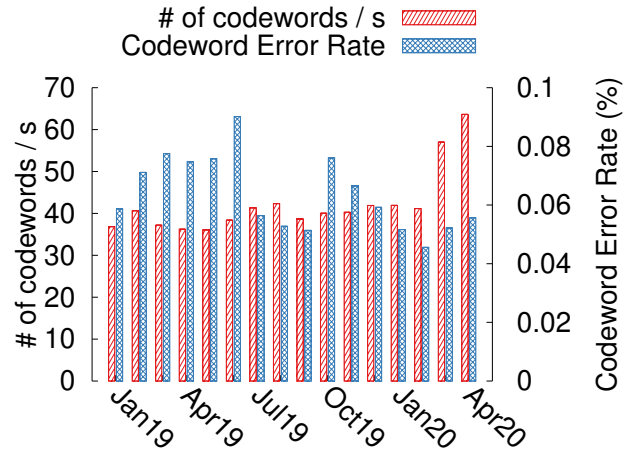


Figure 13: This figure shows the average number of codewords sent per second and the average codeword error rate in each month from Jan 2019 to Apr 2020.

the codeword errors of different devices correlate with each other.

6.1 SNR and TX Power

Figure 11 shows the correlation between the SNR and the codeword error rates. For each data point, we plot the SNR value on the x-axis and the y-axis is the codeword error rate of each upstream channel in log-scale. The average codeword error rate decreases as the SNR increases, indicating that codeword errors are caused by noises breaching into a cable segment. Figure 12 shows the correlation between the transmission (TX) power and the codeword error rate. The average codeword error rate shows a decreasing trend as the TX power increases until 52 dBm. However, the error rates peak when the TX power reaches 52 dBm or 58 dBm, respectively. This is because the modems have reached their maximum TX power settings. Some of them have their maximum TX Power set to 52 dBm, while some modems have it set to 58 dBm or higher. Figure 12 shows that a modem increases its TX power in response to codeword error rates and codeword error rates will spike when a modem cannot overpower the noises in a cable segment.

6.2 Traffic Load

Our codeword dataset includes data collected from January 2019 to April 2020. During the last two months of the data collection period, COVID-19 hit U.S. and remote learning/working started. We break the data points into different months to plot the monthly average number of codewords sent by each device per second and the monthly average codeword error rate. Figure 13 shows the results. We observe that the number of codewords sent per second increases by 44.36% and 61.11% in March 2020 and April 2020, respectively. In contrast, the monthly codeword error rate fluctuates over time, and we do not see a significant increase or decrease in March

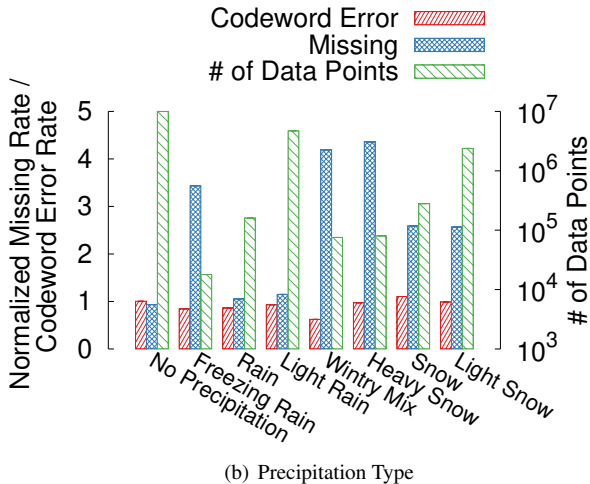
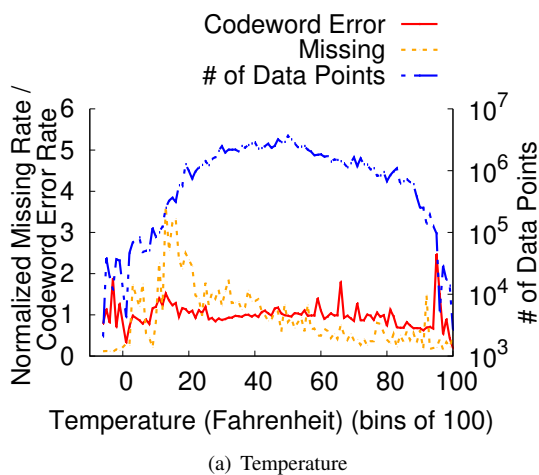


Figure 14: This figure shows how codeword error rates and network (un)availability (captured by the normalized data missing rate) are affected by temperature and type of precipitation.

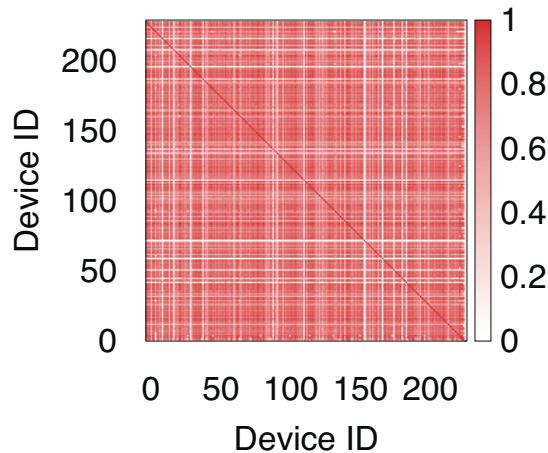


Figure 15: The correlation matrix in the FN with the highest average codeword error rate. Each Device ID represents a modem. This figure shows the codeword error rates of modems in the FN with the highest average codeword error rate are highly correlated.

2020 and April 2020.

6.3 Weather

Padmanabhan et al. [24] have shown that severe weather conditions reduce the availability of residential networks. We are interested in finding out whether severe weather will impact the codeword error rate, which is a network reliability metric. As described in § 3, we collect the historical weather data that overlap in time and location with our codeword data. We compute the codeword error rates under different weather conditions. For comparison, we use the data points where no performance data are collected as indicators of networking being unavailable. We compute the rate when this event happens and refer to it as the missing data rate.

Figure 14(a) shows how the codeword error rate and the

data missing rate change as the temperatures change. For clarity, we normalize the codeword error rate with the average codeword error rate among all FNs seen in our data. Similarly, we normalize the data missing rate. In Figure 14(a), we see the codeword error rate increases when the temperature is around $10^{\circ}F$, $< 0^{\circ}F$, or just below $100^{\circ}F$. We do not have many data points for $< 10^{\circ}F$ or $> 100^{\circ}F$ weather. So the data points in those regions may not be representative. We see that the data missing rate increases significantly when the temperature is between $10^{\circ}F$ and $30^{\circ}F$, consistent with the results in [24].

Figure 14(b) shows the normalized codeword error rate and normalized missing rate in different weather types. The data missing rates are significantly higher in Freezing Rain, Wintry Mix, and Heavy Snow weather types. One explanation

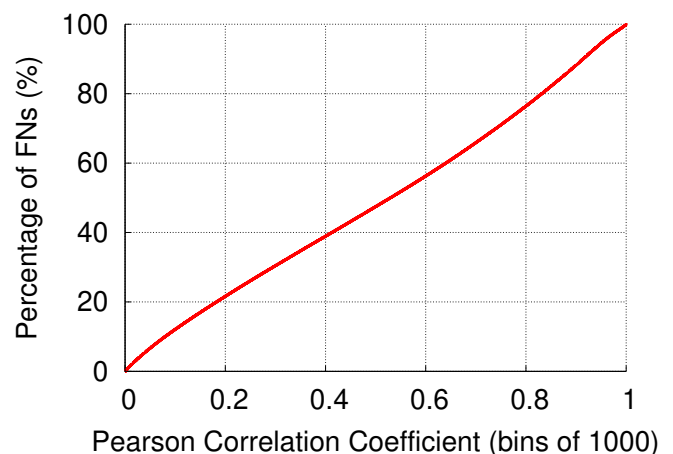


Figure 16: This figure shows the CDF of the average Pearson correlation coefficient of each FN, which is averaged over all devices' pair-wise Pearson correlation coefficients. In about 30% of FNs, the devices show strong error rate correlation (> 0.7).

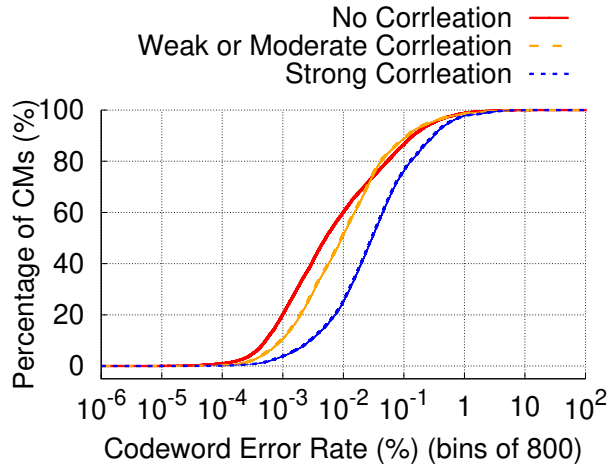


Figure 17: This figure shows the CDF of the average codeword error rates of devices in different correlation groups. The devices that show strong error rate correlation with other devices in the same FN tend to have higher codeword error rates.

we learn from AnonISP is that some HFC networks use aerial cables and these weather types can cause damage to those cables. The weather type’s impact on the codeword error rate is not as clear. The three weather types that significantly increase network unavailability: Freezing Rain, Wintry Mix, and Heavy Snow do not significantly increase the codeword error rate. This indicates the codeword error rate may not be affected by the types of precipitation.

6.4 Codeword Error Correlation

Lastly, we analyze our data to answer this question: *How does a device’s codeword error rate correlate to those of other devices in the same FN?* An HFC network is a shared medium network. Devices connected to the same FN may share RF impairments. Understanding the scope of RF impairment sharing can help us develop future fault diagnosis tools.

We quantify the correlation of codeword error rates of two devices using the pair-wise Pearson correlation coefficient [9]. For each device, we compute its codeword error rate at each data collection point and treat it as an element in the input vector to the Pearson coefficient calculation. Since our data span a 16-month period and each data point is collected every four hours, the length of the vector is around 2.5K. For each FN, we compute the pair-wise Pearson correlation coefficients for all devices in the FN. We then average the Pearson correlation coefficients among all devices in an FN to obtain the average Pearson correlation coefficient of the FN.

Figure 15 shows the correlation coefficient matrix for the FN with the highest average codeword error rate. The average correlation coefficient in this FN is 0.6798. According to [21], a correlation coefficient less than 0.3 indicates no correlation, between 0.3 to 0.7 means weak or moderate correlation, and larger than 0.7 shows a strong correlation. This figure shows that most devices in this FN have a strong error rate correlation. There are also some devices that do not have

any correlation with other devices, suggesting that these devices do not share the same RF impairments with the other devices.

Figure 16 shows the CDF of the average Pearson correlation coefficient in each FN. We observe that the devices in nearly 30% of the FNs show a strong correlation, indicating that for most of the time in these FNs, a large group of devices in the same FN share RF impairments. We also observe that nearly 30% of the FNs present no correlation among the devices.

We are interested in understanding how a device’s codeword error rate is distributed when it shows a certain degree of codeword error rate correlation with other devices in the same FN. We divide the devices into three groups, No Correlation (average Pearson correlation coefficient < 0.3), Weak or Moderate Correlation ($0.3 \leq$ Pearson correlation coefficient < 0.7), and Strong Correlation (Pearson correlation coefficient ≥ 0.7). A device’s average Pearson correlation coefficient is the sum of its pair-wise Pearson correlation coefficients with other devices in the same FN divided by the number of pairs.

Figure 17 shows the CDF of the codeword error rates for devices in different correlation groups. The x-axis is in log-scale. We see that for devices in the Strong Correlation group, the 10th and 90th percentiles of their codeword error rate distributions are $[2.82 \times 10^{-3}\%, 0.306\%]$; for devices in the Weak or Moderate Correlation group, the 10th and 90th percentiles of their codeword error rate distributions are $[9.32 \times 10^{-4}\%, 0.113\%]$; and for devices in the No Correlation group, the 10th and 90th percentiles of their codeword error rate distribution are $[5.21 \times 10^{-4}\%, 0.140\%]$. These results show that devices that show a high correlation to other devices are more likely to have high codeword error rates, suggesting that they are affected by the same RF impairments.

It is possible that an RF impairment only affects a portion of the devices in the same FN, which makes the average Pearson coefficient for the affected devices low since they have no correlation with the unaffected devices. In our future work, we plan to investigate whether a clustering algorithm based on codeword error correlation can help identify the devices that are affected by the same RF impairment.

7 Implications

We believe this work provides several implications for network operations and network research:

- When measuring packet loss on the Internet, one should vary the length of a measurement packet to gain a full spectrum of packet loss statistics. Packets of different sizes may be encoded into different numbers of codewords and experience different loss rates.
- When designing the network applications and protocols, one should take into account packet loss caused by physical-layer transmission errors in the RF systems. Exceptional

innovation in the cable broadband industry has allowed ISPs to use HFC networks to deliver high-speed data. But due to the RF range they operate in, transmission errors in those networks are not negligible.

- ISPs should not rely on customer tickets alone for network maintenance. Customers in chronically high-error-rate networks may have adapted to the network conditions.

8 Related Work

Last-mile Packet Loss: FCC launched the MBA project [2] in 2011 and has been publishing an annual report on broadband performance. FCC MBA project uses UDP pings to measure packet loss. We analyze the FCC data and estimate what fraction of packet loss from the FCC measurements is due to physical-layer transmission errors. Using the FCC data, Sundaresan et al. [28] show that different ISPs and different home network devices can lead to different latency and loss rate distributions. Sundaresan et al. [30] also show that for broadband network customers, the last-mile latency is the main bottleneck when visiting web pages since it significantly contributes to both DNS lookup time and the time to the first byte. Genin and Splett [18] use the download speed distribution from the FCC data to investigate where congestion happens, concluding that most of the Internet congestion occurs in the last-mile network.

In addition to the FCC MBA project, many researchers have also measured and characterized the reliability of the last-mile broadband access networks using their measurement apparatuses. Dischinger et al. [16] measure 1,894 broadband hosts from 11 ISPs with TCP and ICMP measurement packets. They show that both DSL and cable broadband networks exhibit non-negligible packet loss rates, with around 5% data points showing a loss rate higher than 1%. Hu et al. [22] demonstrate that physical layer performance metrics are useful in detecting and predicting network outages that can affect customer experience. Schulman and Spring [26] employ ICMP echo request packets to measure how weather affects the availability of broadband networks. Padmanabhan et al. [25] point out that the last-mile is often the bottleneck by analyzing the end-to-end client-server traffic. Their results indicate that approximately 75% packet loss occur in the last-mile networks. The results presented by Sundaresan et al. [27] support this statement by analyzing the RTT of different TCP traffic. Sundaresan et al. [29] also show the home wireless network is the main bottleneck when a user's access link speed exceeds about 20 Mbps. However, Bajpai et al. [8] measured the last-mile latency in the US and Europe, showing that the last-mile latency is stable over time, which are inconsistent with the observations made by [25, 27, 29]. Fontugne et al. [17] investigate the last-mile latency among 646 ASes, and find that nearly 10% of the ASes presenting persistent last-mile congestion.

Backbone Packet Loss: Apart from the last-mile networks, researchers have also measured the performance of backbone

networks. Ghobadi and Mahajan [19] measure the performance metrics from the optical layer in a large backbone network. Their work shows that one of the optical layer performance metrics, SNR can be used to predict network outages that are not visible to the IP layer. Markopoulou et al. [23] send probes over 43 paths in 7 ISPs to measure the latency and packet loss in the continental US, showing that the packet loss rates for all measured paths are less than 0.26%.

Datacenter Packet Loss: Benson et al. [10] measure packet loss in datacenter networks and show that the packet loss mostly occurs at edge links with low average utilization, indicating the primary cause of packet loss in datacenter networks is momentary spikes. Zhang et al. [31] show most of the packet loss in datacenter networks occur in ToR switches. Both of the studies focus on packet loss in the IP layer. Zhuo et al. [32] show corrupted optical links in datacenter networks introduce a high packet loss rate and the rate of link corruption is not correlated with the link's utilization.

Summary: Different from previous work, this work uses physical-layer codeword statistics to characterize packet loss caused by physical-layer transmission errors. It focuses on the last-mile cable broadband networks and complements previous work.

9 Conclusion

As many applications are sensitive to packet loss, continuously monitoring packet loss in a broadband network has attracted much interest from researchers and policymakers. Previous measurement work, including FCC's decade-long MBA project, cannot differentiate congestion-induced packet loss from transmission-error-induced loss.

This work fills in this blank by using physical-layer data contributed by a cable ISP. The data were collected from 77K+ devices spanning 394 HFC network segments in a 16-month period. Using this data, we infer that physical-layer transmission errors could contribute to more than 12%-25% of packet loss in the cable ISPs measured by the MBA project. We show that some HFC network segments suffer from persistent error loss that exceeds 1%. Customers in these network segments do not make more calls than other customers. These findings suggest that network researchers and operators should take into account packet loss caused by physical-layer errors in network measurement, protocol design, and network maintenance tasks.

Acknowledgment

The authors would like to thank our shepherd Andreas Haeberlen and the anonymous NSDI reviewers for their valuable feedback and the industry experts, especially Jacob Malone, from CableLabs and David Clark for providing insightful suggestions and feedback. Many thanks go to AnonISP for providing us the opportunity to share this work with the research community. This work was supported in part by an NSF award CNS-1910867 and a gift from CableLabs.

References

- [1] Broadband Internet Regulation and Access: Background and Issues. <https://www.everycrsreport.com/reports/RL33542.html>, 2008.
- [2] Measuring Broadband America. <https://www.fcc.gov/general/measuring-broadband-america>, 2011.
- [3] DOCSIS Codeword Errors And Their Effect on RF Impairments. <http://www.zcorum.com/wp-content/uploads/DOCSIS-Codeword-Errors-Their-Effect-on-RF-Impairments.pdf>, 2013.
- [4] Measuring Broadband America (Technical Appendix to the Tenth MBA Report). <https://data.fcc.gov/download/measuring-broadband-america/2020/Technical-Appendix-fixed-2020.pdf>, 2020.
- [5] Measuring Broadband America (Validated Data Cleansing, Tenth Report). <https://data.fcc.gov/download/measuring-broadband-america/2020/validated-data-cleansing-sept2019.pdf>, 2020.
- [6] IBM Environmental Intelligence Suite: Weather Data APIs. <https://www.ibm.com/products/environmental-intelligence-suite/data-packages>, 2021.
- [7] Number of Fixed Broadband Subscribers in the United States from 2010 to 2020. <https://www.statista.com/statistics/217938/number-of-us-broadband-internet-subscribers/#statisticContainer>, 2021.
- [8] Vaibhav Bajpai, Steffie Jacob Eravuchira, and Jürgen Schönwälder. Dissecting Last-mile Latency Characteristics. *ACM SIGCOMM Computer Communication Review*, 47:25–34, 2017.
- [9] Jacob Benesty, Jingdong Chen, Yiteng Huang, and Israel Cohen. Pearson Correlation Coefficient. In *Noise reduction in speech processing*, pages 1–4. Springer, 2009.
- [10] Theophilus Benson, Aditya Akella, and David A Maltz. Network Traffic Characteristics of Data Centers in the Wild. In *ACM IMC*, 2010.
- [11] DOCSIS CableLabs. Best Practices and Guidelines, PNM Best Practices: HFC Networks (DOCSIS 3.0). Technical report, CM-GL-PNMP-V03-160725, 2016.
- [12] DOCSIS CableLabs. Data-Over-Cable Service Interface Specifications DOCSIS® 3.0 Operations Support System Interface Specification. Technical report, CM-SP-OSSIV3.0-C01-171207, 2017.
- [13] DOCSIS CableLabs. Data-Over-Cable Service Interface Specifications DOCSIS® 3.0 Physical Layer Specification. Technical report, CM-SP-PHYv3.0-C01-171207, 2017.
- [14] Neal Cardwell, Yuchung Cheng, C Stephen Gunn, Soheil Hassas Yeganeh, and Van Jacobson. BBR: Congestion-based Congestion Control: Measuring Bottleneck Bandwidth and Round-trip Propagation Time. *Queue*, 14(5):20–53, 2016.
- [15] Wireline Competition. Restoring Internet Freedom. <https://docs.fcc.gov/public/attachments/FCC-17-166A1.pdf>, 2017.
- [16] Marcel Dischinger, Andreas Haeberlen, Krishna P Gummadi, and Stefan Saroiu. Characterizing Residential Broadband Networks. In *ACM IMC*, 2007.
- [17] Romain Fontugne, Anant Shah, and Kenjiro Cho. Persistent Last-mile Congestion: Not so Uncommon. In *ACM IMC*, 2020.
- [18] Daniel Genin and Jolene Splett. Where in the Internet is Congestion? *arXiv preprint arXiv:1307.3696*, 2013.
- [19] Monia Ghobadi and Ratul Mahajan. Optical Layer Failures in A Large Backbone. In *ACM IMC*, 2016.
- [20] Sangtae Ha, Injong Rhee, and Lisong Xu. CUBIC: A New TCP-friendly High-speed TCP Variant. *ACM SIGOPS operating systems review*, 42(5):64–74, 2008.
- [21] Dennis E Hinkle, William Wiersma, and Stephen G Jurs. *Applied Statistics for the Behavioral Sciences*, volume 663. Houghton Mifflin College Division, 2003.
- [22] Jiyao Hu, Zhenyu Zhou, Xiaowei Yang, Jacob Malone, and Jonathan W Williams. CableMon: Improving the Reliability of Cable Broadband Networks via Proactive Network Maintenance. In *USENIX NSDI*, 2020.
- [23] Athina Markopoulou, Fouad Tobagi, and Mansour Karam. Loss and Delay Measurements of Internet Backbones. *Computer communications*, 29:1590–1604, 2006.
- [24] Ramakrishna Padmanabhan, Aaron Schulman, Dave Levin, and Neil Spring. Residential Links Under the Weather. In *ACM SIGCOMM*. 2019.
- [25] Venkata N Padmanabhan, Lili Qiu, and Helen J Wang. Server-based Inference of Internet Link Lossiness. In *IEEE INFOCOM*, 2003.
- [26] Aaron Schulman and Neil Spring. Pingin’ in the Rain. In *ACM IMC*, pages 19–28, 2011.

- [27] Srikanth Sundaresan, Mark Allman, Amogh Dhamdhere, and Kc Claffy. TCP Congestion Signatures. In *ACM IMC*, 2017.
- [28] Srikanth Sundaresan, Walter De Donato, Nick Feamster, Renata Teixeira, Sam Crawford, and Antonio Pescapè. Broadband Internet Performance: A View From the Gateway. 41:134–145, 2011.
- [29] Srikanth Sundaresan, Nick Feamster, and Renata Teixeira. Home Network or Access Link? Locating Last-mile Downstream Throughput Bottlenecks. In *International Conference on Passive and Active Network Measurement*. Springer, 2016.
- [30] Srikanth Sundaresan, Nazanin Magharei, Nick Feamster, and Renata Teixeira. Characterizing and Mitigating Web Performance Bottlenecks in Broadband Access Networks. In *ACM IMC*, 2013.
- [31] Qiao Zhang, Vincent Liu, Hongyi Zeng, and Arvind Krishnamurthy. High-resolution Measurement of Data Center Microbursts. In *ACM IMC*, 2017.
- [32] Danyang Zhuo, Monia Ghobadi, Ratul Mahajan, Klaus-Tycho Förster, Arvind Krishnamurthy, and Thomas Anderson. Understanding and Mitigating Packet Corruption in Data Center Networks. In *ACM SIGCOMM*, 2017.

A Raw Data of Codeword Error Rates

This appendix section includes sample figures of the raw codeword data we used for the analysis in this paper. Figure 18 includes codeword error rates of modems under different conditions. We draw these figures using the data collected from 12 modems between July 1st to July 31st. Figure 18(a) to Figure 18(d) show the data collected from modems in unhealthy FNs. They have high codeword error rates as expected. Figure 18(e) to Figure 18(h) show the data collected from modems in alarming FNs, while Figure 18(i) to Figure 18(l) show the data collected from modems in healthy FNs.

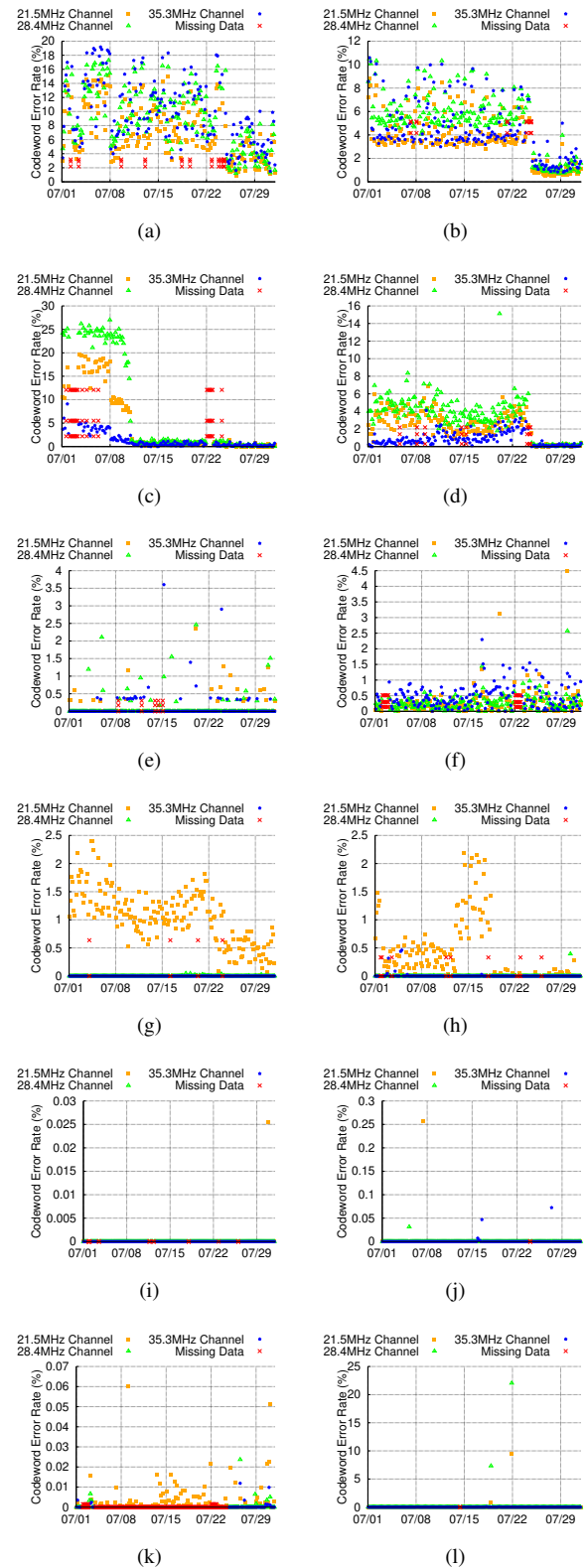


Figure 18: Raw codeword error rates from sample devices.