

A Dynamic Peer-to-Peer Traffic Limiting Policy for ISP Networks

Chaojong Wang, Ning Wang, Michael Howarth
University of Surrey
Guildford, United Kingdom
{C.Wang, N.Wang, M.Howarth}@surrey.ac.uk

George Pavlou
University College London
London, United Kingdom
G.Pavlou@ee.ucl.ac.uk

Abstract—As a scalable paradigm for content distribution at Internet-wide scale, Peer-to-Peer (P2P) technologies have enabled a variety of networked services, such as distributed file-sharing and live video streaming. Most existing P2P systems employ non-intelligent peer selection algorithms for content swarming which greedily consume Internet bandwidth resources. As a result, Internet service providers (ISPs) need some efficient solutions for managing P2P traffic within their own networks. A common practice today is to block or shape P2P traffic in order to conserve bandwidth resources for carrying standard traffic from which revenue can be generated. In this paper, instead of looking at simple time-driven blocking/limiting approaches, we investigate how such types of limiting behaviors can be more gracefully performed by the ISP by taking into account the dynamics of both P2P traffic and of standard Internet traffic. Specifically, our approach is to adaptively limit excessive P2P traffic on critical network links that are prone to congestion, based on periodical link load/utilization measurements by the ISP. The ultimate objective is to guarantee non-P2P service capability while trying to accommodate as much P2P traffic as possible based on the available bandwidth resources. This approach can be regarded as a complementary solution to the recently proposed collaboration-based P2P paradigms such as P4P. Simulation results show that our approach not only eliminates performance degradation of non-P2P services that are caused by overwhelming P2P traffic, but also accommodates P2P traffic efficiently in both existing and future collaboration-based P2P network scenarios.

1 INTRODUCTION

In the past decade, the development of P2P technology has been highly successful in offering a variety of Internet services, ranging from traditional file sharing to real-time multimedia applications such as live video streaming [1, 2] and IP Telephony [3]. On the other hand, the uncontrolled behavior of P2P systems in consuming Internet bandwidth means they account for some 50%-70% of the overall Internet traffic [4, 5]. Under these circumstances, the network's capacity for other services, such as conventional web-based applications may be impacted due to the potential resource competition with overwhelming P2P traffic; in addition, the P2P traffic generally does not create additional revenues for the ISP. This situation has driven many ISPs to seek effective solutions to reduce or even eliminate P2P traffic within their networks. Nevertheless, simply blocking all P2P traffic may not be an appropriate option for ISPs because such decisions may

significantly impact the market reputation of the ISP, given the popularity of P2P services used by their customers today.

As far as blocking/limiting mechanisms on P2P traffic are concerned, Cisco Systems implemented in its commercial products several management policies for controlling P2P traffic, such as aggregated rate limiting [6]. In general, the approach is to *statically* allocate a certain percentage of bandwidth capacity to P2P traffic. However, since traffic patterns in today's ISP networks are highly dynamic even within a single day [7], such an offline approach is unlikely to be efficient in dealing with network dynamics. For example, less P2P traffic might be accommodated despite the availability of unused resources in the network. On the other hand, much research has recently been carried out on improving peer selection behavior at the application layer in order to reduce P2P traffic across the Internet. Most of the proposed schemes have been based on collaboration between the P2P service provider and ISPs in order to optimally utilize network resources [4, 8, 9]. For instance, Aggarwal et al. proposed a generic Oracle service [4] that takes advantage of relevant network layer information when performing locality-based peer selection. This approach allows ISPs and P2P service providers to establish a collaborative relationship in gracefully provisioning P2P services across the Internet. Choffnes et al. [8] implemented an extended BitTorrent system based on the DNS redirection information for selecting peers in an optimized way. In addition, Xie et al. proposed a revolutionary P2P portal architecture called P4P [9] that is able to optimally consume network resources for supporting generic P2P applications. In the P4P framework, the distinct objectives of ISPs and P2P systems are decoupled and realized in a distributed fashion by exchanging P4P-distance that can be defined in various ways such as locality information (e.g., in terms of OSPF link weight) or network status (e.g., congestion information). These proposals aim to reduce or optimize P2P traffic at the application layer via collaboration with ISPs. Nevertheless, it is still to be investigated how ISPs will *independently* take effective control of P2P traffic without relying on collaboration with P2P systems. Research activities have also been carried out in the IETF, with a relevant new working group having been recently formed [7].

In this paper, we introduce a dynamic P2P traffic control policy based on measured traffic that can be independently used by ISPs. Specifically, excessive P2P traffic is blocked

from consuming bandwidth resources on links once their measured utilization exceeds a threshold set by the ISP in advance. In this approach, traditional non-P2P traffic, which effectively generates revenue for ISPs when being carried by the underlying networks, is treated with higher priority compared to P2P traffic. In other words, non-P2P traffic is not actively limited, but P2P traffic is throttled if necessary. Today, many network operators periodically perform network measurements on their traffic volumes and bandwidth utilization, typically with an interval of 5 to 15 minutes [7]. The measured network conditions can be used as input for intelligently limiting P2P traffic in order to best utilize the available bandwidth for both P2P and non-P2P traffic. This proposed dynamic policy enables ISPs to control P2P traffic in a more adaptive way in comparison to the current time-driven approach [6]. In addition, this approach is ISP-centric and it does not assume necessary collaboration between the ISP and P2P systems. From this point of view, the policy can be viewed as a complementary approach to P4P. To implement such an approach, the ISP needs to have mechanisms in place that can identify P2P traffic within their networks. In [2, 10] both an active crawler and a passive technique for identifying P2P flows from other types of traffic were introduced. These approaches are based on specific types of protocols, port numbers, and traffic patterns to identify P2P traffic. Alternatively, P2P traffic can be also identified based on connection patterns without relying on packet payload [11], or through a set of heuristics derived from the robust properties of P2P traffic [12]. Any of these approaches can be used to identify P2P traffic and it is not our intention to focus on a specific one in this paper.

In addition to the introduction of a P2P traffic control policy at the ISP side, we also examine here how individual peers may react to network configuration changes in the application layer; for example, the peer connections may change following the enforcement of limiting P2P traffic somewhere in the network. In other words, we investigate both the ISP network performance and in parallel the P2P service performance, even if there is no collaboration between the two. We focus on an intra-domain network environment (GEANT [13]), since periodically measured traffic volume information is available [13, 14], and this allows us to use a real network topology and background traffic levels in our experiments. In Figure 1 two autonomous entities are shown – the P2P system and the ISP, together with background traffic (i.e. non-P2P traffic). First of all, the P2P system selfishly selects the best partner peers based on its own application layer measurements (for example, end-to-end delay). The ISP then periodically measures the overall network utilization and applies its limiting policy by blocking *excessive* P2P traffic on some critical network links. The ISP’s goal is to protect non-P2P traffic from potential congestion, while accommodating as much P2P traffic as possible where bandwidth is still available. Following the blocking of some network links to excessive P2P traffic, the P2P system reconfigures its peer selection decisions, possibly based on its application-layer measurements. Such configuration changes at the application layer may further impact the network performance measured by the ISP in the next interval. Given such an interaction process, the P2P

system and the ISP can be modeled as two distinct players with different objectives, and they adjust their own strategies without being aware of the explicit behavioral change on the other side. Based on this interaction model, we study the impact of applying our proposed P2P traffic limiting policy on both the ISP’s network condition and the P2P service performance. We find that our approach enables ISPs to control excessive P2P traffic efficiently with appropriate setting of limiting threshold values in both non-collaborative and collaborative P2P environments.

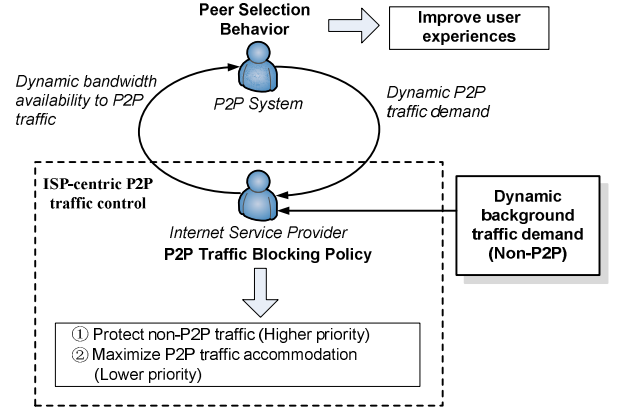


Figure 1: Dynamic Interactions between ISP and P2P

2 SYSTEM MODEL & SPECIFICATION

In this section we first provide the modeling specification for P2P connection sessions (CSs) at the application layer, and underlying network information. The dynamic P2P traffic control scheme based on the proposed limiting policy is then presented. In addition, we analyze how the P2P system reacts to actions dynamically taken by the ISP.

2.1 Physical & P2P networks

We model a physical Point-of-Presence (PoP) network topology as a unidirectional graph $G = (V, E)$, where V is a set of PoP nodes and E is the set of inter-PoP links. In our modeling, each peer in a P2P session is associated with one of the PoP nodes in the physical network topology. According to the common practice of ISP network design, bandwidth resources within a single PoP are usually highly over-provisioned, so we only focus on bandwidth resources on inter-PoP links in E . This means if multiple peering neighbors belong to the same PoP, then the associated bandwidth consumed by their internal peering connections is ignored. Let P_{ij} represent the physical path between PoP nodes i and j , consisting of one or more inter-PoP links. We consider multiple simultaneous P2P connection sessions (CSs) running over the network. Each CS contains a distinct set of active peers sharing the same contents. If one end user participates in multiple CSs, it is treated as an independent peer in different CSs. It should also be noted that peers can only select their own partners for content swarming in each CS but the actual delivery paths are determined by the ISP’s routing configuration, for instance using OSPF/IS-IS.

From the viewpoint of the ISP's physical network, both P2P traffic and conventional non-P2P traffic are carried on top of the same network topology and between each PoP node pair i and j are routed via the same path. We denote such aggregated traffic demand on PoP node pair ij as

$$t_{ij} = t_{ij}^p + t_{ij}^{np} \quad (1)$$

where t_{ij}^p and t_{ij}^{np} denote the overall P2P traffic demand and the overall conventional non-P2P traffic demand from PoP node i to j respectively.

A traffic demand t_{ij} is delivered on a distinct path P_{ij} . We introduce a binary mapping coefficient Y_{ij}^l to indicate the relationship between t_{ij} and a physical link l : Y_{ij}^l equals 1 if physical link l constitutes path P_{ij} ($l \in P_{ij}$) and hence serve t_{ij} , and 0 otherwise. Therefore, the total traffic demands T_l on the physical link l can be expressed as

$$T_l = \sum_{i \in V} \sum_{j \in V, i \neq j} t_{ij} * Y_{ij}^l, Y_{ij}^l = \begin{cases} 1 & \text{if } l \in P_{ij} \\ 0 & \text{otherwise} \end{cases} \quad (2)$$

At the P2P system side, let V^p denote a set of active peers in the network. Each requesting peer ($u \in V^p$) needs to download contents from its partners ($V_u \subset V^p$) in the same CS at a certain transmission rate. Usually, a requesting peer is known as being fully served in a CS, provided that its aggregated content swarming rate with all other partners has reached or exceeded a certain minimum level. For instance, a fully served peer in a real-time P2P video streaming application might be able to play back the stream at a stable quality based on a specific data rate, e.g., 1Mbps [1, 2]. In our model we thus use D_u to denote a given transmission rate at which a full served requesting peer u is supposed to download in a particular peering session. We define a peering session of a requesting peer u to be a set of connections between u and all its partners (identified by set V_u) in a specific CS. Due to the fact that the uplink bandwidth of each peer (defined as UB_{vu} from partner v to requesting peer u) is usually lower than the overall demand of a requesting peer in today's ASDL environments, we therefore define S_u to be the actual content downloading rate of a requesting peer u from its partners in V_u as

$$S_u = \sum_{u \neq v, \forall v \in V_u} UB_{vu} \quad (3)$$

If $S_u \geq D_u$, then the requesting peer u is regarded as a fully served peer.

2.2 Dynamic P2P traffic limiting policy

As we have mentioned, due to the highly dynamic behavior of P2P and non-P2P (background) traffic, pure static or time-driven P2P traffic control policies may not make efficient use of bandwidth resources. For instance, all P2P traffic might be blocked even if there is available bandwidth, or on the other hand, excessive P2P traffic might be admitted into the network upon a sporadic upsurge in the number of P2P connections, leading to service deterioration of non-P2P flows. In order to solve these problems, we propose a dynamic P2P traffic limiting policy that can be deployed at the ISP side. As we have indicated, current network measurements are

typically performed at short intervals, for instance between 5 to 15 minutes, and hence it is possible that our approach of controlling P2P traffic at a relatively short timescale can be applied in such environments.

Figure 2 illustrates the workflow for our dynamic P2P traffic limiting policy algorithm. First, the ISP measures the current network load on per inter-PoP link basis, allowing the actual link utilization to be computed. Using a *pre-determined link utilization threshold* for limiting P2P traffic, the ISP then determines any link on which excessive P2P traffic needs to be blocked during the current interval. Specifically, when the utilizations of all links in the network are below the threshold set by the ISP, all P2P traffic can traverse these links. If the measured utilization has exceeded the threshold on some links, some P2P traffic on those links will be blocked in order to bring down the overall utilization in the current interval. To do this, a proportion of P2P traffic is randomly selected for such blocking, for instance based on source/destination address prefixes. As a result, connections with some existing peers may be affected, and hence re-selection of partner peers at the application layer is required. In addition, peers that have newly arrived during the current interval are also unable to select the potential partner peers whose associated paths contain the P2P-limited links. From the description above, we can find that our policy allows both ongoing and new P2P sessions to consume bandwidth resources until a certain level is reached. Even when the overall link utilization exceeds the threshold, only some P2P traffic is blocked in order to guarantee high service priority for non-P2P traffic. It should be noted that the scope of such traffic limiting is only performed based on the original measurement in the beginning of each interval. Note that even if some links that were not originally impacted by the limiting action (because their original measured utilization was lower than the threshold) have increased utilization exceeding the threshold due to *peer re-selections* and *newly joined peers* in the middle of the current interval, there is no new rate limiting on such links. Hence, congestion may still occur within each interval between two adjacent bandwidth limiting actions. Nevertheless, through careful tuning of the threshold, such a possibility can be effectively minimized. The rationale behind this strategy is that blocking some excessive P2P connections in a proactive way on some highly utilized links will force the impacted peers to reselect their partners, which will lead to some P2P traffic being diverted away from those links. As a result, P2P traffic can be more optimally redistributed across the network, and more incoming traffic can be accommodated during the current interval. It is also worth mentioning that, under certain circumstances where the network has already been highly loaded, some peers may fail to identify a sufficient number of partner peers for supporting the ongoing sessions (ref to Eq.3) due to P2P wide-scope traffic limiting actions – this is particularly the case in real-time P2P applications which require a minimum level of downloading rate.

To formally specify our policy, we recall that E is the set of inter-PoP links, and a new variable E'_r ($E'_r \subseteq E$) is defined

as a set of *unblocked* inter-PoP links that can be used for carrying P2P traffic in interval τ without any limiting constraint. We also define the utilization function of each physical link l as $U_\tau(l) = T_l / C_l$ where C_l is the capacity of physical link l . Let θ_τ be the utilization threshold of links regulated by the ISP at interval τ . Figure 2 outlines how to determine links for limiting P2P traffic as well as the corresponding reactions at the P2P system side. As an ISP-centric approach, such a scheme can be applied not only to the current non-collaborative approaches between P2P systems and the underlying ISPs, but also to emerging paradigms that foresee possible collaboration between the two players, for instance the P4P portal service.

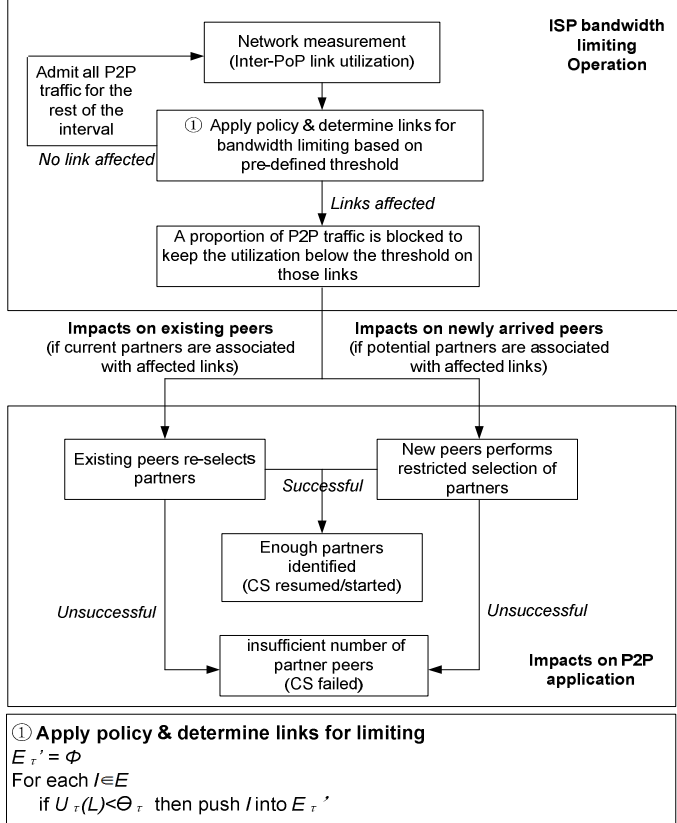


Figure 2: Dynamic P2P Traffic Limiting Policy

3 EVALUATION METHODOLOGY

In this section we describe our simulation environment, including the network topology, (background) traffic demands, and peer dynamics patterns. In addition, we also describe the performance metrics used in our evaluations.

3.1 Network Topology

We use the real GEANT network topology [13] at the PoP-level, which consists of 23 nodes and 74 unidirectional inter-PoP links. Each link has its actual link capacity and IGP link weight. According to [7], we know that the IGP link weight setting is based on end-to-end latency in the GEANT network, and hence customer traffic is effectively routed on the lowest delay paths.

3.2 Traffic Demands

Our simulation is based on the GEANT network traffic traces over 24 hours. According to [7], the GEANT traffic traces are measured every 15 minutes through NetFlow. In our evaluation we take the samples of these traces at 15-minute interval during the period of 24 hours. Figure 3 shows the measured maximum link utilization (MLU) performance in the GEANT network across these 24 hours (starting from 00:00 am), indicated with 96 intervals each of 15 minutes. It can be clearly seen that the overall traffic volume is highly dynamic, for instance the minimum value of MLU during the period is around 35%, while the maximum value in the same day can reach as high as 85%. In our evaluation, we use the scaled volume of these traffic traces to emulate the non-P2P background traffic behavior.

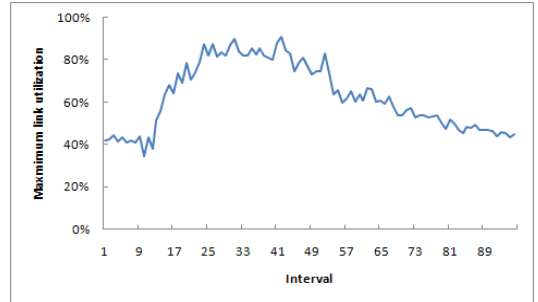


Figure 3: Background Traffic Dynamics in GEANT

The P2P traffic used in our experiment is synthetically generated according to the flow characteristics of today's popular real-time multimedia based P2P applications [1, 2]. We consider 20 concurrent P2P channel sessions, with each channel attracting up to 1200 peers. Hence altogether we consider up to 24000 peers that are distributed across the 23 PoP nodes in the GEANT network. The overall distribution of these peers in each PoP node is determined according to the population of each city, where larger PoP nodes have more peers assigned. The channel session selected by each peer is randomly determined. In addition, we use the observation that each requesting peer has around 30 peering connections in order to satisfy the overall downloading rate requirement for playback in a stable peering session (1Mbps, [1, 2]). In each peer connection there is one top peer which contributes on average three times contents as others based on the measurement of a popular real-time P2P content delivery system [2].

3.3 Peer Join & Departure Patterns

We used a popular probability model to determine P2P session dynamics (i.e. if an event is a peer join or leaving), which has been proposed in [15]. The function is:

$$P_x(k_x) = \frac{\alpha_x(n_x - k_x)}{\alpha_x(n_x - k_x) + (1 - \alpha_x)k_x}$$

Here, P_x is the probability used for event generation – whether a new event is a peer join or leaving in the current connection session x . The metric of k_x is defined as the number of active peers in the current connection session x , and n_x is the maximum number of peers in the connection session. In addition, the variable α_x lying in the range of (0, 1) is used to

control the popularity of connection session x . On the other hand, the authors of [2, 16] measured and captured a variety of participating peers in a single P2P-based IPTV channel. Based on their measurements we performed some moderate adaptations on the above probability modeling equation in order to emulate more close to the actual P2P dynamics pattern. Figures 4 and 5 indicate how the total number of active peers varies across all the sessions under consideration during the period of one single day. The dynamics of total number of participating peers follow similar pattern as indicated in [2] which is based on the PPLive application. Figure 5 shows the actual peer join and departure patterns.

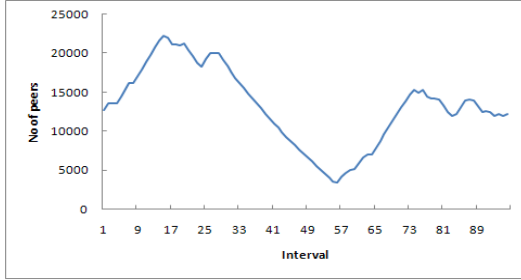


Figure 4: Total no. of active peers in P2P system

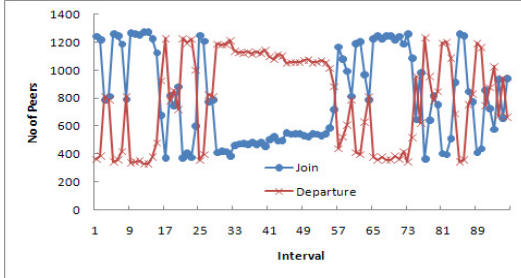


Figure 5: Peer arrival & departure patterns

3.4 Evaluation Metrics

We evaluate both P2P system performance and network performance in our experiments. On the P2P system side, we define the following two metrics - *partner peer churn ratio* and the *insufficiency ratio* to indicate how the dynamic P2P traffic limiting policy impacts the P2P system performance at the application layer.

- *Partner peer churn ratio (CR)*: Due to the unavailability of some network links following the bandwidth limiting actions at the ISP side, a set of existing active peers may need to re-select some of their partners to replace the original ones whose connections are impacted by the limiting action. This scenario may lead to perceivable performance disruptions since it needs to take time to re-select and re-connect new partners. This metric is defined as:

$$CR = \frac{\text{No. of Requesting Peers In-Churn}}{\text{Total No. of Requesting Peers}}$$

- *Insufficiency ratio (IR)*: In each peering session, the requesting peer is supposed to download the contents from its partners at a minimum transmission rate (see equation 3). By using this metric we can capture the total number of such

requesting peers that cannot be fully served following the limiting actions by the ISP.

$$IR = \frac{\text{No. of None Fully Served Requesting Peers}}{\text{Total No. of Requesting Peers}}$$

In addition, network performances are evaluated according to the following two metrics.

- *Maximum link utilization (MLU)*: MLU is defined across the entire PoP-level network topology, and is often used to indicate potential network congestion levels. More specifically, the lower the maximum link utilization is, the less chance that traffic congestion in the network will occur. MLU is defined as:

$$MLU = \max_{l \in E} (U_l) = \max_{l \in E} \left(\frac{T_l}{C_l} \right)$$

- *Network cost*: The piece-wise linear cost function has been widely used for evaluating traffic engineering purposes. In this paper we use the cost function proposed in [17], that is:

$$\varphi = \sum_{l \in E} \varphi_l(U_l)$$

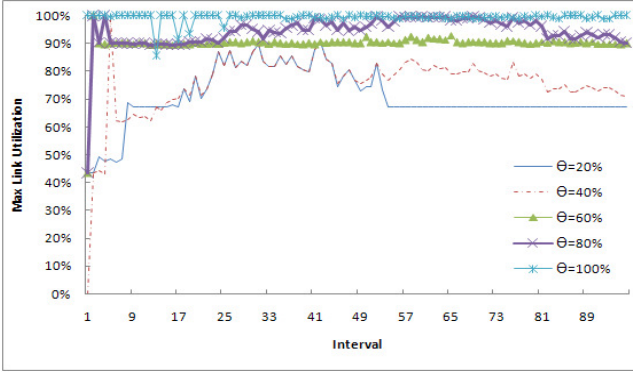
Where for all $l \in E$, $\varphi_l(0) = 0$ and

$$\varphi_l(x) = \begin{cases} 1 & \text{for } 0 < x \leq 1/3 \\ 3 & \text{for } 1/3 < x \leq 2/3 \\ 10 & \text{for } 2/3 < x \leq 9/10 \\ 70 & \text{for } 9/10 < x \leq 1 \\ 500 & \text{for } 1 < x \leq 11/10 \\ 5000 & \text{for } 11/10 < x \leq \infty \end{cases}$$

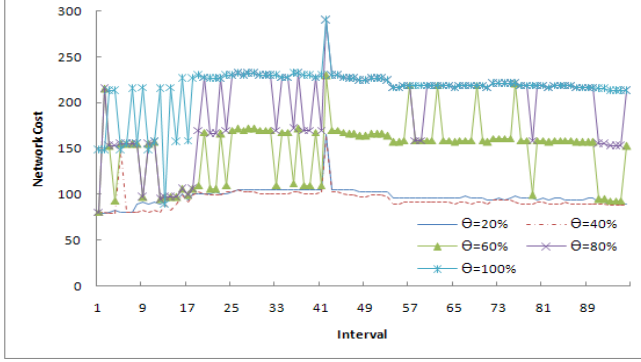
As far as the configuration of threshold θ is concerned, we set a range of limiting thresholds at 20%, 40%, 60%, 80% and 100% in our evaluation. It can be easily inferred that the lower the threshold θ is, the more “conservative” the ISP’s policy is in accommodating P2P traffic. Since bandwidth limiting actions are taken every 15 minutes, peers join between two adjacent blocking actions may incur additional P2P traffic, and this can be also be accompanied by an increase of background non-P2P traffic. As a result, some network links may still suffer from congestion, and packets of both P2P and background traffic may get dropped as a consequence. This scenario is regarded as conventional router-incurred traffic loss, in comparison to active packet drop behavior based on an instantaneous bandwidth limiting policy.

4 PERFORMANCE EVALUATION

We now evaluate the performance at both the network and the P2P system. We first focus on the non-collaborative environment where no information is exchanged between ISP and P2P system, with peers being selected in a greedy fashion without taking into account network conditions. We then consider a collaborative P2P model such as P4P where peer selection decisions are made according to the network conditions provided from the underlying ISP network. To achieve this we implemented a simplified P4P-like system based on the specification in [9].

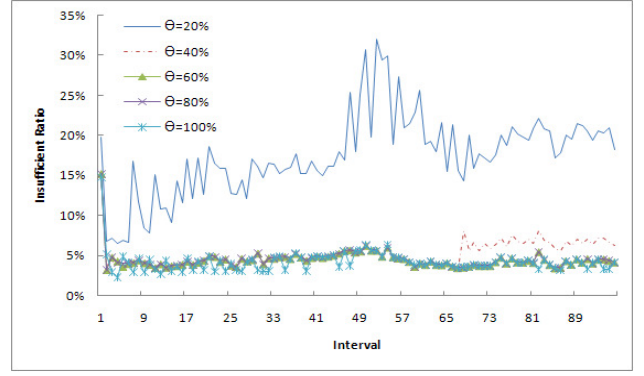


(a) Max Link Utilization (MLU)

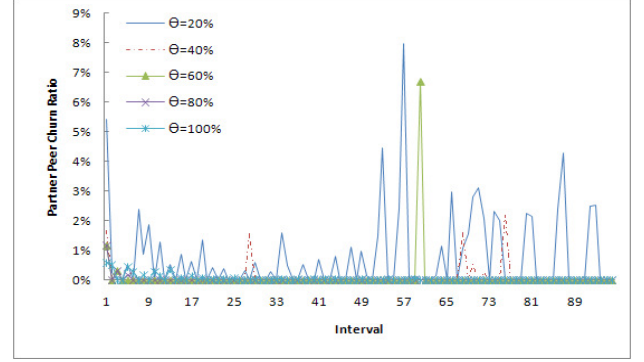


(b) Network Cost

Figure 6: Network Performances



(a) Insufficiency Ratio



(b) Partner Peer Churn Ratio

Figure 7: P2P System Performances

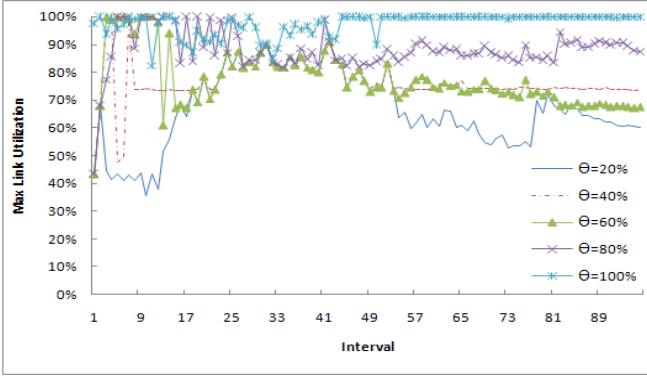
4.1 Non-collaborative P2P environments

We first focus on the non-collaboration scenario in which the P2P system greedily selects partner peers with the objective of minimizing end-to-end delay, but without taking into account the underlying network conditions. The ISP employs our dynamic P2P traffic limiting policy based on link utilization to block excessive P2P traffic. This case can be regarded as the existing P2P network scenario in today's Internet. We study the P2P system performance and the network performance under different threshold settings, and analyze how the tuning of such metric will impact the performance on both sides.

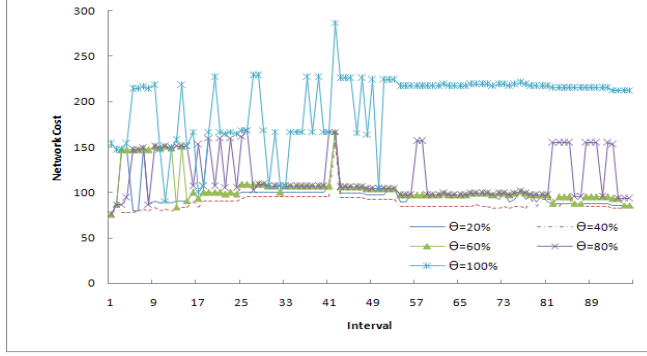
Threshold $\theta=20\%$: Figure 6 shows the overall network performances (in terms of both MLU and network cost). When the utilization threshold is set to 20% (i.e. extra P2P traffic is blocked on any link whose measured overall link utilization exceeds 20%), the MLU performance curve somehow follows a similar pattern to Figure 3, which is the actual background traffic in the GEANT network. This observation is expected since most of P2P traffic is blocked unless the background traffic volume becomes sufficiently low. This observation of low MLU is also echoed in Figure 6 (b) where $\theta=20\%$ has a low network cost. On the P2P performance side (Figure 7), with $\theta=20\%$, we can see that the corresponding insufficiency ratio is by far the highest among all the threshold configurations, as shown in Figure 7(a). More specifically, around 15% of the requesting peers are not able to be fully served. In fact, we observed that all surviving requesting peers are actually connecting mostly to partners located in the same

PoP, i.e. where we assume high available bandwidth links. When the local partners are not able to provide the minimum content swarming rate, then the requesting peer needs to seek partners from outside its local PoP. In the worst case, if still no feasible partners can be found, for instance due to the widely blocked network links, the requesting peer will become starved. A similar situation happens to the churn ratio shown in Figure 7 (b), which indicates highly unstable P2P session connections. In summary, this conservative threshold setting achieves good network performance at the ISP side, but these are at the expense of relatively poor P2P system performance at the user side.

Threshold $\theta=40\%$, 60% , and 80% : we analyze the performance under the limiting threshold at 40%, 60% and 80% together. From Figure 6 (a), we can see that the 40% threshold corresponds to the lowest MLU performance among these three scenarios. In addition the 40% threshold also achieves lower network cost than 60% and 80% thresholds (Figure 6 (b)). In fact its network cost performance is very close to the 20% scenario. As for the churn ratio performance at the P2P side, the 40% threshold increases by around 10% on average in comparison to the other two thresholds at several intervals. By comparing the performance of these three thresholds, we observe that the insufficiency and churn ratios are roughly on the same level, but the 40% threshold has noticeably lower MLU. This indicates that more background traffic can be potentially accommodated with such setting but no significant performance degradation on the P2P side will be incurred.

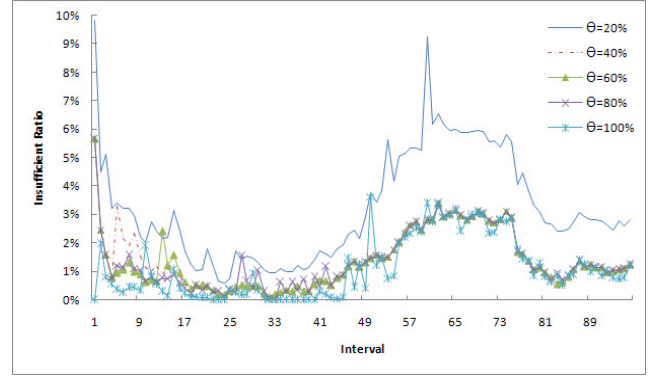


(a) Max Link Utilization (MLU)

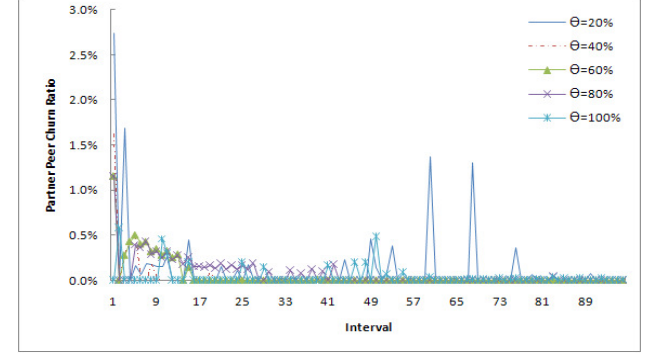


(b) Network Cost

Figure 8: Network Performance



(a) Insufficiency Ratio



(b) Partner Peer Churn Ratio

Figure 9: P2P System Performance

Threshold $\theta=100\%$: A special case is to set the limiting threshold at 100%. In this scenario there is no bandwidth limiting on P2P traffic until 100% link utilization is detected. This configuration can be regarded as a pure “reactive” approach in which the network may have already experienced congestion before the actual measurement. We can clearly see from Figure 6 (a) that 100% threshold results in 100% utilization for most intervals, whilst, as expected, it also incurs the highest network cost.

4.2 Collaborative P2P Environments

Collaborative P2P paradigms have recently been proposed in the literature which allow more intelligent peer selections based on network conditions. In this paper we choose the P4P architecture [9] as an example. We have implemented a simplified P4P simulation platform based on [9]. According to the P4P architecture, the P2P system is able to make use of the underlying information on the network condition provided by ISP to optimally select partner peers. Whilst it may take some time to deploy such paradigms widely across the Internet, we nevertheless study the performance of our threshold configuration policy on such advanced P2P platforms.

Threshold $\theta=20\%$: From Figure 8 (a) we can see that the performance pattern under threshold 20% again follows almost the same as the background traffic pattern shown in Figure 3, since most of the P2P traffic experiences some “early” blocking. In Figure 8 (b) the performance curve of 20% threshold overlaps with that of the 60% threshold in network cost performance. On the P2P system side, as shown in the performance of insufficiency and churn ratios (Figure 9),

we can see that there is an average 82% decrease in comparison to the non-collaborative case since P4P is able to reduce the MLU significantly thanks to the intelligent peer selections with network condition awareness. In the worst case across all intervals, around 10% of all requesting peers suffer from insufficiency, and the churn ratio is kept under 3% (as compared to 32% and 8% respectively in the non-collaborative scenario). By comparing the corresponding performance between Figures 6/7 and 8/9, we can see that the setting of the threshold based traffic blocking to the P4P paradigm results in even better performance on the P2P system, since partner peers are optimally selected based on underlying network information for avoiding certain high utilization links.

Threshold $\theta=40\%$, 60%, and 80%: Compared with the more conservative 20% threshold, 40% and 60% thresholds have very similar MLU performance with each other as shown in Figure 8 (a), while the 80% threshold produces higher MLU than other two. In Figure 8 (b), the network cost performance of the 40% threshold is slightly lower than that of 60% and 80%, and even 20%. This observation is consistent with the non-collaborative case so that we consider the 40% threshold to be an appropriate configuration. In Figure 9 (a) for the insufficiency performance, the curves for 40%, 60% and 80% thresholds overlap with each other for most of the period, so we consider that they have the same insufficient performance, and this is lower by 82% on average than the non-collaboration case. We also find that 40%, 60%, 80% thresholds have much lower churn ratios in the network towards the end of the measurement period.

Threshold $\theta=100\%$: The final case is the 100% threshold configuration. From Figure 8 (a) we can see that this threshold gives a 100% actual MLU during most of the period, and this high traffic loading is also reflected by the highest network cost in Figure 8 (b). On the other hand, the insufficiency ratio curve in Figure 9 (a) is almost the same as that of the 40% and 60% thresholds, while the 100% threshold has more fluctuations compared with the others. In addition, the churn ratio of the 100% threshold stays lower, close to the 40% and 60% threshold settings of Figure 9 (b). Therefore the 100% threshold is not the most appropriate configuration since it has the highest MLU (100%) despite of similar churn and insufficient ratios to the 40% and 60% thresholds.

4.3 Key observations

Based on our experiments in the *non-collaborative* P2P environment, we can see that the 40% threshold scenario seems to reach a promising trade-off between both sides – the maximum link utilization is relatively lower, and in addition no significant performance degradation is observed on the P2P system side. On the other hand, the 40% and 60% threshold settings have lower insufficiency ratio, churn ratio, and better MLU performance in the *collaborative* P2P environment. These proper threshold settings indicate that more background traffic can be potentially accommodated without significant performance degradation on P2P systems. This observation is consistent between the non-collaboration and the collaboration scenarios.

Based on the above observations, we conclude that 1). P2P traffic control actions should be taken in a proactive manner, rather than waiting until network congestion has been actually detected (e.g., the $\theta=100\%$ case). 2). Both too conservative and too greedy limiting threshold settings for accommodating P2P traffic may result in suboptimal performance on either side. Careful tuning of this threshold based on both network engineering objectives and service requirements from P2P users is essential. 3). There exists at least an appropriate configuration of threshold setting for bandwidth limiting, possibly leading to a “win-win” situation in both non-collaborative and collaborative P2P paradigms. However the threshold setting should not be unique across various network and traffic patterns. Nevertheless it does provide the possibility for individual ISPs to tweak such a value according to their own situations and objectives. 4). Collaborative-based P2P systems such as P4P can benefit even more from our approach in terms of both network performance and application efficiency. From this point of view, such a P2P traffic blocking policy can be regarded as a complementary solution to these emerging collaboration paradigms.

5 CONCLUSIONS AND FUTURE WORK

It has been observed that current P2P traffic management policies adopted by ISPs (whenever this is the case) are not intelligent enough and not able to cope efficiently with highly dynamic traffic patterns and increasing volumes of P2P traffic. Collaborative P2P architectures such as P4P aim to improve the situation at the application level through information provided by the underlying ISPs. In this paper we propose instead a dynamic P2P traffic limiting policy that can be independently applied by ISPs in order to control P2P traffic to meet their

own service/operational objectives; for instance, to provide higher treatment priority to standard non-P2P traffic while accommodating as much P2P traffic as possible by taking into account dynamic bandwidth availability. Since our approach is ISP-centric, it can be applied to both non-collaboration based and collaboration based P2P networks. In our performance evaluation based on the GEANT network, we discovered that certain “optimal” configuration of threshold values exist, being able to achieve global optimization for the performance of both the underlying network and the overlaid P2P systems.

In our future work, we will first extend our experiments to evaluate additional service performance metrics, such as the end-to-end delay at the P2P user level, which can be also affected by dynamic bandwidth limiting actions. Moreover, we also plan to investigate the behavioral interaction between P2P systems and more complicated traffic engineering (TE) paradigms rather than simple traffic blocking approaches. Such analysis can be based on game-theory approach, e.g. [18, 19].

6 ACKNOWLEDGMENT

This work was partially funded by the EU FP7 COMET Project (248784).

7 REFERENCES

- [1] SopCast, <http://www.sopcast.com>
- [2] X. Hei et al., “A measurement study of a large-scale P2P IPTV system”, IEEE Trans on Multimedia, Vol. 9, No. 8, pp. 1672-1687, Dec. 2007
- [3] Skype, <http://www.skype.com>
- [4] V. Aggarwal et al., “Can ISPs and P2P Users Cooperate for Improved Performance?”, ACM CCR, Vol. 37, No. 3, pp. 29-40, July 2007
- [5] The IETF ALTO WG: <http://www.ietf.org/dyn/wg/charter/alto-charter.html>
- [6] White Paper, “Management Peer-to-Peer traffic with Cisco service Control Technology”, http://www.cisco.com/en/US/prod/collateral/ps7045/ps6129/ps6133/ps6150/prod_white_paper0900aecd8023500d.html
- [7] U. Steve et al., “Providing Public Intradomain Traffic Matrices to the Research Community”, ACM CCR, Vol. 36, No. 1, pp. 83-86, Jan. 2006.
- [8] D.R. Choffnes et al., “Taming the Torrent: A Practical Approach to Reducing Cross-ISP Traffic in P2P Systems”, ACM CCR, Vol. 38, No. 4, pp. 363-374, Oct. 2008
- [9] H. Xie et al., “P4P: Provider Portal for Applications”, Proc. ACM SIGCOMM 2008
- [10] S. Sen et al., “Analyzing Peer-To-Peer Traffic Across Large Network”, IEEE Trans. on Networking, Vol. 12, No. 2, pp. 219-232, April 2004
- [11] K. Thomas et al., “Transport Layer Identification of P2P Traffic”, Proc. ACM IMC 2004
- [12] P. Marcell et al., “Identification and Analysis of Peer-to-Peer Traffic”, Journal of Communications, Vol. 1, No. 7, pp. 36-46, Nov. 2006
- [13] The GEANT Network topology, <http://www.geant.net>
- [14] N. Spring et al., “Measuring ISP topologies with Rocketfuel”, IEEE/ACM Trans on Networking, Vol. 12, No. 1, pp. 2-16, Feb. 2004
- [15] B.M. Waxman, “Routing of Multipoint Connections”, IEEE Journal on Selected Areas in Communications, Vol. 6, No. 9, pp. 1617-22, Dec. 1988
- [16] L. Vu et al., “Measurement and Modeling of a Large-scale Overlay for Multimedia Streaming”, Proc. ACM Qshine 2007
- [17] B. Fortz et al., “Internet Traffic Engineering by Optimizing OSPF Weights”, Proc. IEEE INFOCOM 2000.
- [18] L. Qiu et al., “On Selfish Routing in Internet-Like Environments” Proc. ACM SIGCOMM 2003
- [19] Y. Liu et al., “On the Interaction Between Overlay Routing and Underlay Routing”, Proc. IEEE INFOCOMM 2005