

Exploring the Trade-off Between Performance and Data Freshness in Database-Driven Web Servers

Alexandros Labrinidis

Department of Computer Science
University of Pittsburgh
labrinid@cs.pitt.edu

Nick Roussopoulos

Department of Computer Science
University of Maryland
nick@cs.umd.edu

Abstract

Personalization, advertising, and the sheer volume of online data generate a staggering amount of dynamic web content. In addition to web caching, View Materialization has been shown to accelerate the generation of dynamic web content. View materialization is an attractive solution as it decouples the serving of access requests from the handling of updates. In the context of the Web, selecting which views to materialize must be decided online and needs to consider both performance and data freshness, which we refer to as the Online View Selection problem. In this paper, we define data freshness metrics, provide an adaptive algorithm for the online view selection problem which is based on user-specified data freshness requirements, and present experimental results. Furthermore, we examine alternative metrics for data freshness and extend our proposed algorithm to handle multiple users and alternative definitions of data freshness.

1 Introduction

The frustration of broken links from the early Web has been replaced today by the frustration of web servers stalling or crashing under the heavy load of dynamic content. In addition to data-rich online web services, even seemingly static web pages are usually generated dynamically in order to include personalization and advertising features. However, dynamic content has significantly higher resource demands than static web pages (at least one order of magnitude) and creates a huge scalability problem at web servers.

Dynamic web caching [IC97, CIW⁺00, DDT⁺01, DDT⁺02, LKM⁺02, ABK⁺03] has been proposed to solve this scalability issue. The biggest problem of employing caching techniques for dynamic web content is the coupling of serving access requests and handling of updates, since an update that invalidates a cached object will result in the object being recomputed on the next access request. For example, imagine a cache that can store dynamically generated web pages. During normal operation we get an 80% hit rate (which means that only 20% of the pages will need to be recomputed). If we get a small surge in the update stream, a big percentage of the cached pages could be invalidated, and the hit rate will drop significantly. A sudden drop in hit rate leads to a sudden increase in the average response time and possibly to server saturation.

View materialization can solve this problem, since it decouples the serving of access requests from the handling of the updates.

Selecting which views to materialize, the *view selection problem*, has been studied extensively in the context of data warehouses[Rou82, Gup97, GM99, MRSR01]. However, unlike data warehouses, which are off-line during updates, most web servers maintain their back-end databases online and perform updates concurrently with user accesses. Therefore, in the context of the Web, selecting which views to materialize must be decided dynamically and needs to consider both performance and data freshness.

In this paper, we present $\text{OVIS}(\theta)$, an adaptive algorithm for the Online View Selection problem. $\text{OVIS}(\theta)$ acts as a knob in the system, determining at run-time which views should be materialized (cached and re-freshed immediately on updates) and which ones should just be cached and re-used as necessary. Parameter θ corresponds to the level of data freshness that is considered acceptable for the current application. In addition to maintaining high performance given the data freshness demands, $\text{OVIS}(\theta)$ also detects infeasible thresholds, when the freshness demands would create a backlog at the web server.

Motivating Example: Our motivating example is a database-driven web server that provides realtime stock information to subscribers. Updates to stock prices and other market derivatives are streamed to the back-end database and must be performed online. The web server provides users with up-to-date information that includes current stock prices, moving average graphs, comparison charts between different stocks and personalized stock portfolio summaries. In general, we are interested in data-intensive web servers that provide mostly dynamically generated web pages to users (with data drawn from a DBMS) and also face a significant online update workload.

Structure of paper: In the next section, we present our metrics for measuring system performance and data freshness. We also define the Online View Selection Problem. In Section 3 we describe the proposed Online View Selection Algorithm and in Section 4 we discuss the results of our experiments. Section 5 examines alternative metrics for data freshness, whereas Section 6 extends the proposed Online View Selection Algorithm to handle multiple users and alternative definitions of data freshness. Section 7 summarizes related work. We conclude in Section 8.

2 Problem Definition

We extend the typical three-tier architecture of modern web servers, by adding an *Asynchronous Cache* module, between the application server and the database server (Figure 1). In this architecture, the web

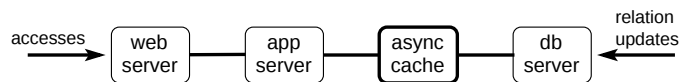


Figure 1: System Architecture

server module is responsible for serving user requests and the application server is responsible for web workflow management. Instead of interfacing the application server directly to the database server, the Asynchronous Cache module acts as an intermediary. Unlike traditional caches in which data is simply

invalidated on updates, data in the asynchronous cache can be *materialized* and immediately refreshed on updates. Recent products from IBM and Microsoft incorporate such an asynchronous middle-tier cache [BAK⁺03, LGZ03, LGZ04].

2.1 Web Page Derivation Graph

There are three types of data objects in the system: relations, WebViews, and web pages.

- **Relations** are stored in the database server and are the primary “storage” for structured data. They are affected by the incoming update stream. Relation updates are executed in order of arrival.
- **WebViews**, introduced in [LR00], are HTML or XML fragments. In other words, WebViews are simply parts of a Web page. WebViews are usually generated by “wrapping” database query results (i.e. database views) with HTML formatting commands or XML semantic tags. We allow WebViews to be formed from *any type* of database queries. In fact, the only assumptions that we must make for WebViews are that:
 1. we have their current definition and corresponding query (which means that we must have the values of any parameters if the WebViews are generated by a parameterized query)
 2. are able to determine when they become stale (in the most conservative case, this corresponds to whenever the source relations are updated), and
 3. can uniquely address/name them (since we plan to cache them and must be able to perform lookups).

We prefer the term *WebView* over the term “HTML fragment”, which was introduced earlier, in order to stress that these HTML fragments are derived from a database. In fact, we will use the terms “view” and “WebView” interchangeably for the rest of the paper. In the general case, WebViews may be defined from other WebViews, but in this paper we focus on HTML WebViews that are directly derived from querying the database (as we will see later).

- **Web pages** are composed of one or more WebViews. Web pages are what the user is served with in response to his/her access requests.

A *WebView* W_j is derived from relation R_i if W_j includes data which is generated by querying R_i . A web page P_k is considered to be derived from *WebView* W_j if P_k contains W_j .

The associations between these data objects are depicted using a *Web Page Derivation Graph*, which is a directed acyclic graph. The nodes of the graph correspond to relations, WebViews, or web pages. An edge from node a to node b exists only if node b is derived directly from node a . A node can have multiple “parents”, therefore the in-degree of a node can be bigger than one. Relations are the roots of the graph, with zero in-degree, and web pages are the leafs of the graph, with zero out-degree. Figure 2 has an example of a Web Page Derivation Graph. We assume a database with three relations (R_1, R_2, R_3), four WebViews (W_1, W_2, W_3, W_4) and two web pages (P_1, P_2).

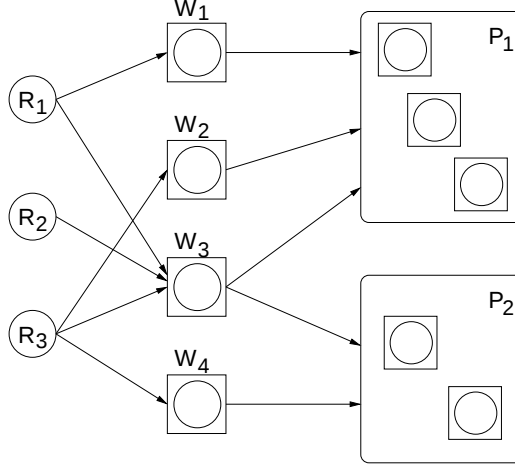


Figure 2: Web Page Derivation Graph

Figure 2 is a very small example of an actual Web Page Derivation Graph. In practice, we usually have thousands of web pages in a web site, with dozens of HTML/XML fragments on each page [CIW⁺00]. However, we also expect to have a significant amount of WebView “sharing” among these web pages. Imagine, for example, a personalized newspaper site. Each user selects the type of news to be included (e.g. local, national, economy), specifies a city for the weather forecast, and gives a list of stock symbols along with the purchase price and quantities for calculating his/her portfolio value. Although the combination of the above elements is most probably unique, there is clearly a finite number of cities/stock symbols, which will be shared among thousands of users (in addition to the standard navigation/presentation fragments).

2.2 The Asynchronous Cache

All requests that require dynamically generated content are intercepted by the *Asynchronous Cache module* (ASC for short). ASC maintains WebViews using one of the following three policies.

Virtual WebViews are always executed on demand and never cached. Intercepted queries against Virtual WebViews are forwarded to the database server, whereas database updates do not affect them.

Non-Materialized (cached) WebViews are cached in ASC, in anticipation of future requests. While they are fresh, they are served very efficiently from the cache. When an update affects a WebView, the cached WebView is invalidated and needs to be re-generated on a following request. This is similar to traditional caching with invalidation rather than a Time-to-Live (TTL) consistency protocol. Assuming that invalidating “dirty” WebViews in ASC is not a costly operation, non-Materialized WebViews is always a better policy than Virtual, since, without any loss in data freshness, one obtains significant improvement in response time (for the times when a fresh version of the WebView is in the ASC). Serving from ASC results in two orders of magnitude improvement in response time compared to querying the DBMS.

Materialized WebViews are cached and continuously maintained in the presence of updates. Accesses to them are *always served from the ASC*. The response time is similar to a fresh non-materialized WebView. We assume that the response time remains almost constant, since a Materialized WebView is served from

ASC even when it is not fresh. However, there is a limit as to how many WebViews should be materialized. Materializing too many WebViews increases the overhead of refreshing them all in the background and can have a negative effect on both server performance and WebView freshness.

The big difference between materialized and non-materialized WebViews is the **decoupling** of serving access requests from updating WebViews. With materialization, updates are not in the critical path of serving user requests. Without materialization, updates must be taken care of while serving user requests (i.e. by refreshing a stale WebView before responding). This decoupling helps materialized WebViews avoid increases in response times when there is an increase in update volume (and thus an increase in the percentage of invalidated cached WebViews). In addition to providing data storage, the asynchronous cache module is responsible for automatically selecting which WebViews to materialize.

In this work we consider HTML WebViews only, since we expect them to have the most impact (being more widely deployed than XML WebViews). Dealing only with HTML WebViews means that the cost to generate any WebView from other WebViews will be negligible (simple concatenation of HTML fragments), compared to the cost of generating the WebViews from relational data. This would be in contrast for example with XML WebViews, which may have complicated transformations and derivations from other (parent) XML WebViews. Such an environment becomes very similar to that of the traditional view materialization problem, where one must consider the view derivation dependencies [Rou82, Gup97, GM99, LSTV99, TLS01, MRSR01]

Since the cost to generate WebViews from other WebViews is negligible, in this work we only consider materializing WebViews which are generated directly from relational data (stored in the database server) and do not consider materializing WebViews derived from other WebViews. As was suggested in the literature [LR00, YFIV00], response times for WebViews generated from relational data can be reduced dramatically if they are materialized. Thus, they are the only ones that could offset the overhead of materialization (keeping them up to date in the background). Finally, we assume that WebViews are refreshed by recomputation. This is the general case, which assumes that there is no method for incrementally refreshing the materialized WebViews.

2.3 Measuring Performance

We define the performance of data-intensive web servers by observing the incoming access request stream for a time interval T and measuring the average response times for each user request.

Definition 1: *Performance* is measured as the average response time for user requests.

Specifically, we measure the time between the arrival of the request at the web server and the departure of the response. We measure response times at the web server, since all our techniques aim at improving the performance of the web server.

Improving web server performance might actually not be visible to the end user. Even a ten-fold improvement in response time at the server (e.g., from 100 msec to 10msec) can stay undetected by end users who will receive their responses after a few seconds of network delay. However, a ten-fold performance improvement at the web server clearly improves **scalability**: the same web server configuration can serve

ten times more users or handle sudden ten-fold surges in traffic without the cost of additional hardware.

2.4 Measuring Quality of Data

The “goodness” of the results generated by data-intensive web servers has been neglected in the past. However, with web servers being used for increasingly important applications (e.g. stock market information), it is crucial to measure and improve the *Quality of Data* served to the users. One common characteristic across data-intensive web servers is their *online nature*: updates to source data are applied concurrently with user accesses, since web servers are always available and never off-line. Therefore, the freshness of data served is the most important measure of Quality of Data.

Definition 2: *Quality of Data (QoD)* for data-intensive web servers is the average freshness of the served web pages.

When an update to a relation is received, the relation and all data objects that are derived from it become *stale*. Database objects remain stale until an updated version of them is ready to be served to the user.

We illustrate this with an example. Let us assume the Web Page Derivation Graph of Figure 2, and that only WebViews W_1 and W_2 are materialized. If an update on relation R_1 arrives at time t_1 , then relation R_1 will be stale until time $t_2 \geq t_1$, the time when the update on R_1 is completed (Figure 3). Although we do not cache relations, relation R_1 will be considered stale because of the unapplied update during $[t_1, t_2]$. On the other hand, materialized WebView W_1 will be stale from time t_1 until time $t_3 \geq t_2$, when its refresh is completed. If an update on relation R_3 arrives at a later time, t_4 , then relation R_3 will be stale for the $[t_4, t_5]$ time interval, until t_5 , when the update on R_3 is completed (Figure 3). Also, non-materialized WebViews W_3 and W_4 will be stale for the same interval $[t_4, t_5]$. On the other hand, materialized WebView W_2 will be stale from time t_4 until time $t_6 \geq t_5$, when its refresh is completed.

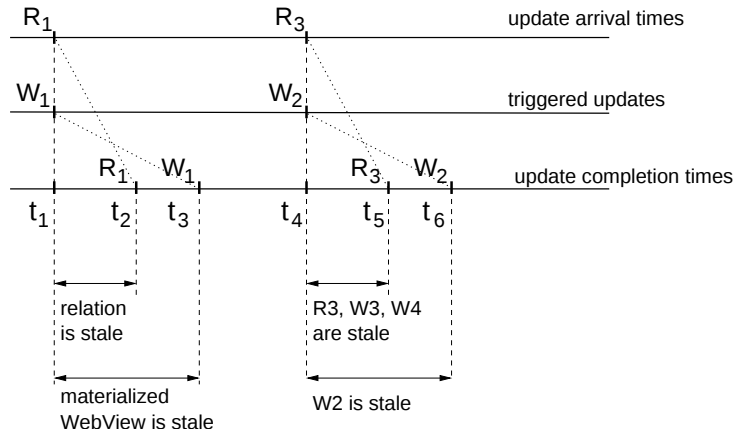


Figure 3: Staleness Example

We identify four types of data objects that can be stale: relations, non-materialized WebViews, materialized WebViews, and web pages.

- **Relations** are stale when an update for them has arrived, but not yet executed.

- **Non-materialized WebViews** are stale when an update for a parent relation has arrived, but not yet executed.
- **Materialized WebViews** are stale if the WebViews have not been refreshed yet (after an update to a parent relation).
- **Web pages** are stale if a parent WebView is stale.

In order to measure freshness, we observe the access request stream and the update stream for a certain time interval T . We view the access stream during interval T as a sequence of n access requests:

$$\dots, A_x, A_{x+1}, A_{x+2}, \dots, A_{x+n-1}, \dots$$

Access requests A_x are encoded as pairs (P_j, t_x) , where t_x is the arrival time of the request for web page P_j . Each web page P_j consists of multiple HTML fragments (WebViews).

We define the freshness function for a WebView W_i at time t_k as follows:

$$f(W_i, t_k) = \begin{cases} 1, & \text{if } W_i \text{ is fresh at time } t_k \\ 0, & \text{if } W_i \text{ is stale at time } t_k \end{cases} \quad (1)$$

A WebView W_i is *stale*, if W_i is materialized and has been invalidated, or if W_i is not materialized and there exists a pending update for a parent relation of W_i . A WebView W_i is *fresh*, otherwise.

In order to quantify the freshness of individual access requests, we recognize that web pages are based on multiple WebViews. A simple way to determine freshness is by requiring that all WebViews of a web page be fresh in order for the web page to be fresh. Under this scheme, even if one WebView is stale, the entire web page will be marked as stale. In most occasions, a strict Boolean treatment of web page freshness like this will be inappropriate. For example, a personalized newspaper page with stock information and weather information should not be considered completely stale if all the stock prices are up to date, but the temperature reading is a few minutes stale.

Since a strict Boolean treatment of web page freshness is impractical, we adopt a *proportional definition*. Web page freshness is a rational number between 0 and 1, with 0 being completely stale and 1 being completely fresh. To calculate $f(A_k)$, the freshness value of web page P_j returned by access request $A_k = (P_j, t_k)$ at time t_k , we take the weighted sum of the freshness values of the WebViews that compose the web page:

$$f(A_k) = f(P_j, t_k) = \sum_{i=1}^{n_j} a_{i,j} \times f(W_i, t_k) \quad (2)$$

where n_j is the number of WebViews in page P_j , and $a_{i,j}$ is a weight factor.

Weight factors $a_{i,j}$ are defined for each (WebView, web page) combination and are used to quantify the *importance* of different WebViews within the same web page. Weight factors for the same web page must sum up to 1, or $\sum_{i=1}^{n_j} a_{i,j} = 1$, for each web page P_j . When a WebView W_i is not part of web page P_j , then the corresponding weight factor is zero, or $a_{i,j} = 0$. The user does not have to specify weight factors. By default, these are set to $a_{i,j} = \frac{1}{n_j}$, where n_j is the number of WebViews in page P_j (which gives all

WebViews equal importance within the same page). However, weight factors can also be user-defined, to reflect bigger importance of a WebView (fragment) within a page, compared to the other WebViews in the same page. Such definitions are always page-dependent.

The overall Quality of Data for the stream of n access requests will then be:

$$QoD = \frac{1}{n} \times \sum_{k=x}^{x+n-1} f(A_k) \quad (3)$$

2.5 Online View Selection Problem

The choice of WebViews to materialize will have a big impact on performance and data freshness. On the one extreme, materializing all WebViews will give high performance, but can have low quality of data (i.e. views will be served very fast, but can be stale). On the other hand, keeping all views non-materialized will give high quality of data, but low performance (i.e. views will be as fresh as possible, but the response time will be high).

We define the *Online View Selection problem* as follows: in the presence of continuous access and update streams, dynamically select which WebViews to materialize, so that overall system performance is maximized, while the freshness of the served data (QoD) is maintained at an acceptable level. In addition to the incoming access/update streams, we assume that we are given a web page derivation graph (like the one in Figure 2), and the costs to access/update each relation/WebView.

Given the definition of QoD from Section 2.4, an acceptable level of freshness will be a threshold $\theta \in [0, 1]$. For example, a threshold value of 0.9 will mean that roughly 90% of the accesses must be served with fresh data (or that all web pages served are composed of about 90% fresh WebViews).

The view selection problem is characterized *online* for two reasons. First, since updates are performed *online*, concurrently with accesses, we must consider the freshness of the served data (QoD) in addition to performance. Second, since the accesses and updates are continuously streaming into the system, any algorithm that provides a solution to the view selection problem must decide at run-time and have the ability to adapt under changing workloads. Off-line view selection cannot match the wide variations of web workloads.

We will use the term *materialization plan* to refer to any solution to the Online View Selection problem. We do not consider the virtual policy for WebViews, since caching will always give as fresh data as the virtual policy and will reuse results, giving better performance. In this paper we assume that the Asynchronous Cache module has infinite size and thus there is no need for a cache replacement algorithm (which would distort the comparison).

To visualize the solution space for the Online View Selection Problem we enumerate all possible materialization plans for a small workload and compute the performance and QoD in Figure 4. The different materialization plans provide big variations in performance and Quality of Data. For example, plans in the bottom left corner of Figure 4 correspond to materializing most WebViews (with very low average response time and low QoD), whereas plans in the top of the plot correspond to not materializing most WebViews (with very high QoD and high average response times).

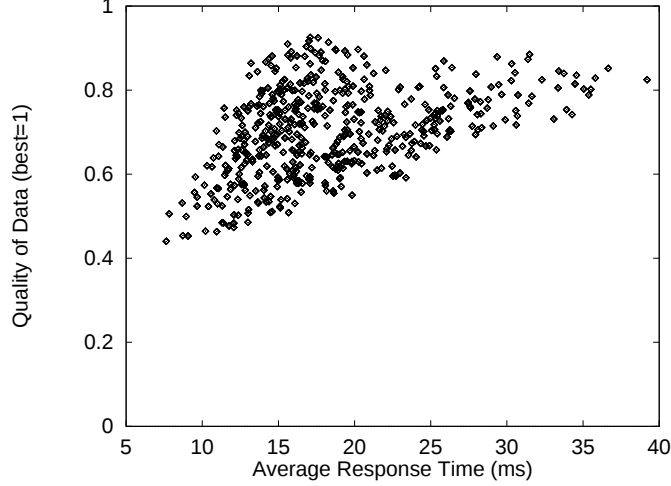


Figure 4: Perf/QoD of all Materialization Plans

3 The OVIS Algorithm

Traditional view selection algorithms work off-line and assume knowledge of the entire access and update stream. Such algorithms will not work in an online environment, since the selection algorithm must decide the materialization plan in real-time. Furthermore, updates in an online environment occur concurrently with accesses, which makes the freshness of the served data an important issue. Finally, the unpredictable nature of web workloads mandates that the online view selection algorithm be *adaptive* in order to evolve under changing web access and update patterns.

In this section we describe $OVIS(\theta)$, an Online View Selection algorithm, where θ is a user-specified QoD threshold. $OVIS(\theta)$ strives to maintain the overall QoD above the user-specified threshold θ and also keep the average response time as low as possible. OVIS also monitors the access stream in order to prevent server backlog.

$OVIS(\theta)$ is inherently adaptive. The algorithm operates in two modes: passive and active. While in passive mode, the algorithm collects statistics on the current access stream and receives feedback for the observed QoD. Periodically, the algorithm goes into active mode, where it will decide if the current materialization plan must change and how.

Figure 5 illustrates the main idea behind the $OVIS(\theta)$ algorithm. By constantly monitoring the QoD for the served data, the algorithm distinguishes between two cases when it must change the materialization plan. When the observed QoD is higher than the threshold θ , $OVIS(\theta)$ identifies a *QoD surplus*, which chooses to “invest” in order to improve the average response time. On the other hand, when the observed QoD is less than the threshold θ , the algorithm identifies a *QoD deficit*, for which it must compensate.

3.1 $OVIS(\theta)$ Statistics

We want to be able to estimate the change in average response time and overall QoD after adapting the materialization plan. We also want to accurately observe the QoD for the served data in order to determine

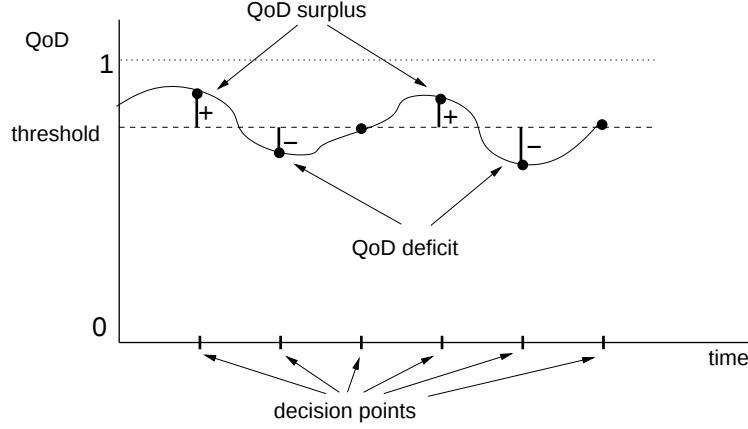


Figure 5: OVIS(θ) Algorithm

whether we have a QoD surplus or a QoD deficit. For that purpose we maintain statistics for each WebView and use them to estimate future behavior. Specifically, we estimate:

- the access frequency for each WebView,
- the performance contribution for each WebView in case it will be materialized and in case it will not be materialized,
- the overall data freshness (QoD) contribution of each WebView in case it will be materialized and in case it will not be materialized, and,
- the amount of change in performance and QoD (differentials) if we change the materialization policy for a WebView.

We explain these statistics, along with the estimation methods, in the next paragraphs.

Estimating the access frequency

The most important statistic in our system is the number of accesses each WebView gets. We must consider popularity, because the materialization decision for popular WebViews will have great impact on both the average response time and the overall QoD. We maintain the total number of accesses for a WebView W_i , which we write as $N_{acc}(W_i)$. The $N_{acc}(W_i)$ counter is incremented whenever there is an access request for a web page that contains W_i .

We use the *Recursive Prediction Error Method* [Jac88] to estimate the number of accesses a WebView will have in the near future. According to this method, we use the measurement for the current period, m , and the previous estimate a , to generate a new estimate a' using the following formula:

$$a' = (1 - g)a + gm \quad (4)$$

where g is a gain factor, $0 < g < 1$. Gain was set to 0.25 for all of our experiments (as was suggested by [Jac88]). As illustrated in Figure 5, the OVIS(θ) algorithm is executed *periodically* in order to adapt the

materialization plan. Periods can be defined either by the number of web page requests received (e.g., adapt every 1000 page requests) or by time intervals (e.g., adapt every 2 minutes). Before each adaptation, we consolidate all statistics and generate estimates for the future. Using Equation 4, we have

$$N'_{acc} = (1 - g)N_{acc} + gN_{acc}^m \quad (5)$$

where N'_{acc} is the new estimate for the number of accesses, N_{acc} is the old estimate for the number of accesses, and N_{acc}^m is the number of accesses measured for the current interval.

Estimating Performance

We estimate the overall cost for implementing a materialization policy and use it to quantitatively compare the changes in performance. High cost will correspond to high average response times and thus low performance.

If a WebView W_i is not materialized, then the overall cost will depend on the *Asynchronous Cache hit ratio*, or how many times we have a cache miss versus a cache hit. Cache misses mandate recomputation of W_i , whereas cache hits will lower the overall cost. We use “ $W_i \not\rightarrow \text{mat}$ ” to denote that W_i will not be materialized. If H_r is the estimate of the hit ratio for WebView W_i , we have:

$$\text{cost}(W_i \not\rightarrow \text{mat}) = \underbrace{H_r \times N_{acc} \times A_{hit}}_{\text{cache hits}} + \underbrace{(1 - H_r) \times N_{acc} \times A_{miss}}_{\text{cache misses}} \quad (6)$$

where N_{acc} is the estimate for the number of accesses for W_i , A_{hit} is the access cost for a cache hit on W_i , and A_{miss} is the access cost for a cache miss on W_i . All estimates are computed using Equation 4. For readability, we do not use the W_i subscripts whenever they can be easily inferred.

The hit ratio, H_r , depends on the materialization policy. If a WebView is materialized, we expect a high hit ratio, because WebViews are refreshed immediately after an update. On the other hand, if a WebView is not materialized, we expect a lower hit ratio (even if eventually the user receives fresh results after cache misses). For that purpose, we maintain separate statistics depending on whether the WebView was materialized or not. When we are trying to estimate the overall cost for a WebView that *will not be* materialized, we use the statistics from when the WebView *was not* materialized. When we are trying to estimate the overall cost for a WebView that *will be* materialized, we use the statistics from when the WebView *was* materialized. The only exception to this is the estimation for the number of accesses and the number of updates which do not depend on the materialization policy.

The hit ratio used in Equation 6 is based on statistics from when WebView W_i was not materialized. If such statistics are not available (because W_i was always materialized in the past), then we use an optimistic estimate for the hit ratio, $H_r = 100\%$.

If a WebView W_i is materialized, the overall cost will not depend on the Asynchronous Cache hit ratio (since all accesses are served from the Asynchronous Cache), but it will depend on the update rate. Updates lead to immediate refreshes and thus impose a computational “burden” on the system. We use $W_i \rightarrow \text{mat}$ to

denote that W_i will be materialized. The overall cost in this case will be:

$$\text{cost}(W_i \rightsquigarrow \text{mat}) = \underbrace{N_{acc} \times A_{hit}}_{\text{accesses}} + \underbrace{R_r \times N_{upd} \times U_{mat}}_{\text{refreshes}} \quad (7)$$

where R_r is an estimate of what percentage of source updates leads to WebView refreshes for W_i , N_{upd} is the estimate of the number of source updates that affect W_i , and U_{mat} is the cost to refresh WebView W_i . The refresh ratio, R_r , is not always 100% because sometimes refreshes are “batched” together (e.g., when there is an update surge). Finally, all estimates are computed using Equation 4.

Eq. 7 assumes that the cost of refreshing the materialized WebViews in the asynchronous cache will impact the response time of serving access requests. This is true when all three software components (Web Server, Asynchronous Cache, DBMS) reside in the same machine, which is a typical configuration for data-intensive web servers today [LKM⁺02].

Estimating the QoD

Similarly to performance, we use statistics to estimate the overall Quality of Data after adapting the materialization plan. Let us assume that $N_{fresh}(W_i | P_j)$ is the number of fresh accesses to WebView W_i which originated from requests to page P_j . The overall QoD definition from Equation 3 can be rewritten as follows:

$$QoD = \frac{1}{n} \times \sum_i \sum_j [a_{i,j} \times N_{fresh}(W_i | P_j)]$$

for all WebViews W_i and all web pages P_j , where n is the total number of page access requests and $a_{i,j}$ are the weight factors defined in Section 2.4. Weights $a_{i,j}$ sum up to 1.0 for all WebViews in the same web page.

Instead of separate $N_{fresh}(W_i | P_j)$ counters for all (WebView, page) combinations, we maintain only one *weighted* counter, $N_{fresh-a}(W_i)$ for each WebView W_i . We increment $N_{fresh-a}(W_i)$ by the weight value $a_{i,j}$ for each fresh access to W_i originating from a request to page P_j . We have that $N_{fresh-a}(W_i) = \sum_j [a_{i,j} \times N_{fresh}(W_i | P_j)]$, for all web pages P_j . Therefore, the QoD definition can be simplified as:

$$QoD = \frac{1}{n} \times \sum_i N_{fresh-a}(W_i) \quad (8)$$

To estimate the contribution of an individual WebView W_i to the overall QoD, we maintain a *freshness ratio*, F_r , defined as $\frac{N_{fresh-a}}{N_{acc-a}}$, where N_{acc-a} is the N_{acc} counter computed similarly to $N_{fresh-a}$ for each WebView W_i . The difference is that N_{acc-a} is incremented by $a_{i,j}$ on every access, not just the accesses that produced fresh results, which is the case for $N_{fresh-a}$. The freshness ratio depends on the materialization policy, therefore we need to maintain separate statistics for when the WebView was materialized and for when it was not materialized, similarly to the hit ratio estimation in the previous section. Given the freshness ratio, F_r , and Equation 8, the QoD contribution for each WebView W_i will be:

$$QoD(W_i) = F_r \times \frac{N_{acc-a}}{n} \quad (9)$$

where n is the total number of web page requests.

Estimating the performance/QoD differentials

At each adaptation step, $\text{OVIS}(\theta)$ must decide if changing the materialization policy for a particular WebView is warranted or not. In other words, it must determine whether it should stop materializing a materialized WebView, or whether it should begin materializing a WebView that was not previously materialized.

After estimating the performance and QoD for all WebViews using the formulas from the previous paragraphs, we compute the performance and QoD differentials for switching materialization policies. For example, if a WebView W_i is currently materialized, we compute the difference in performance and QoD, if W_i were to stop being materialized.

To estimate Δ_{perf} , the **performance differential** for WebView W_i , we use the cost formulas from Eq. 6 and Eq. 7. If W_i is materialized, then we want to estimate how much the performance will change if W_i stops being materialized:

$$\Delta_{\text{perf}} = \text{cost}(W_i \not\rightarrow \text{mat}) - \text{cost}(W_i \rightarrow \text{mat}) \quad (10)$$

Similarly, if W_i is not currently materialized, then we want to estimate how much the performance will change if W_i starts being materialized:

$$\Delta_{\text{perf}} = \text{cost}(W_i \rightarrow \text{mat}) - \text{cost}(W_i \not\rightarrow \text{mat}) \quad (11)$$

A positive performance differential means that the average response time will increase, whereas a negative performance differential means that the average response time will decrease (which is an improvement).

To estimate Δ_{QoD} , the **QoD differential** for WebView W_i , we use the QoD formulas from Eq. 9. If W_i is materialized, then we want to estimate how much the QoD will change if W_i stops being materialized:

$$\Delta_{\text{QoD}} = \text{QoD}(W_i \not\rightarrow \text{mat}) - \text{QoD}(W_i \rightarrow \text{mat}) \quad (12)$$

Similarly, if W_i is not currently materialized, then we want to estimate how much the QoD will change if W_i starts being materialized:

$$\Delta_{\text{QoD}} = \text{QoD}(W_i \rightarrow \text{mat}) - \text{QoD}(W_i \not\rightarrow \text{mat}) \quad (13)$$

A positive QoD differential means that the QoD will increase, (which is an improvement), whereas a negative QoD differential means that the QoD will decrease.

3.2 OVIS(θ) Algorithm

The $\text{OVIS}(\theta)$ algorithm constantly monitors the QoD of the served data and periodically adjusts the materialization plan (i.e. which WebViews are materialized and which ones are not materialized). By maintaining the statistics presented in the previous subsection, $\text{OVIS}(\theta)$ has a very good estimate of how big an effect on

the overall performance and QoD the changes in the materialization plan will have. As we outlined at the beginning of this section, $OVIS(\theta)$ “invests” QoD surplus or tries to compensate for QoD deficit (Figure 5). In the following paragraphs we present the details of the $OVIS(\theta)$ algorithm for the case of QoD surplus and QoD deficit. We also explain why we need to impose a constraint on the maximum amount of change to the materialization plan in a single adaptation step, and, describe how to detect server lag. Finally, we provide the pseudo-code for the $OVIS(\theta)$ algorithm.

QoD Surplus

When the observed QoD Q is higher than the user-specified threshold θ , the algorithm will “invest” the surplus QoD ($= Q - \theta$) in order to decrease the average response time. This is achieved by materializing WebViews which were not previously materialized. For the algorithm to take the most profitable decision, we just need to maximize the total performance benefit, $\sum \Delta_{perf}$, for the WebViews that become materialized, while the estimated QoD “losses”, $\sum \Delta_{QoD}$, remain less than $Q - \theta$. A greedy strategy, that picks WebViews based on their Δ_{perf} improvement provides a good solution, as we explain later.

QoD Deficit

When the observed QoD Q is less than the threshold θ , the algorithm will have to compensate for the QoD deficit ($= \theta - Q$). In this case, $OVIS(\theta)$ will stop materializing WebViews thus increasing QoD, at the expense of increasing the average response time. For the algorithm to take the most profitable decision, we just need to maximize the total QoD benefit, $\sum \Delta_{QoD}$, for the WebViews that stop being materialized, while the estimated overall QoD does not increase above the threshold θ . A greedy strategy, that picks WebViews based on their Δ_{QoD} benefit provides a good solution, as we explain later.

Maximum Change Constraint

Allowing any number of WebViews to change materialization policy during a single adaptation step of the $OVIS(\theta)$ algorithm can have detrimental effects. Since we do not have prior knowledge of the future, any estimate of future performance and QoD after a materialization policy change is just an estimate and can be wrong. Therefore it is preferable to take smaller adaptation “steps”, which should result in a more stable algorithm. For this reason, we impose a limit on the number of WebViews that can change materialization policy during a single adaptation step. We specify this limit as a percentage over the total number of WebViews in the system and denote it as MAX_CHANGE . For example, if $MAX_CHANGE = 5\%$ and we have 1000 WebViews in our system, then at most 50 of them can change materialization policy at a single adaptation step of the $OVIS(\theta)$ algorithm.

Greedy Strategy

With the maximum limit in mind, the desired behavior for $OVIS(\theta)$ under *QoD surplus* can be summarized as follows:

- maximize the improvement in performance, and
- minimize the decrease in QoD

while

- changing the materialization policy of at most MAX_CHANGE WebViews, and
- $QoD > \theta$.

A knapsack-style greedy algorithm (i.e. pick the WebViews with the highest Δ_{perf} per QoD unit) would be preferable if there was no limit to the number of WebViews. However, with the maximum change constraint, a greedy algorithm selecting the top MAX_CHANGE WebViews with the highest Δ_{perf} is the best solution.

Let us see why. In the general case, we assume that because of the MAX_CHANGE constraint, we will not be able to reach our goal of $QoD = \theta$ in a single step. However, we would like to change the materialization policy, so that we get as close as possible to that goal. For this reason, we want to maximize the overall improvement in a single step. For example, if we have 1000 WebViews, and MAX_CHANGE is 10, then we need to identify the 10 WebViews that would give us the highest overall improvement in a single step. Clearly, a greedy strategy that selects the 10 WebViews with the highest Δ_{perf} is the best solution.

Similarly, the desired behavior for $OVIS(\theta)$ under *QoD deficit* can be summarized as follows:

- maximize the improvement in QoD, and
- minimize the decrease in performance,

while

- changing the materialization policy of at most MAX_CHANGE WebViews, and
- $QoD \leq \theta$.

With the maximum change constraint, a greedy algorithm selecting the top MAX_CHANGE WebViews with the highest Δ_{QoD} is the best solution.

Server Lag Detection

From elementary queueing theory [Jai91] we know that system performance worsens dramatically as we approach 100% utilization. In practice, there can be cases where the incoming access and update workload generate more load than what the server can handle, resulting in backlog, which we refer to as *server lag*. Figure 6 has an example of server lag, which is visible on the response times of the access requests.

It is crucial to detect server lag in an online system. For users, server lag means near-infinite response times - this holds for both current (i.e. those still waiting a response) and future users of the system. For system administrators, failure to identify server lag can lead to long, ever-increasing backlogs which will eventually crash the server.

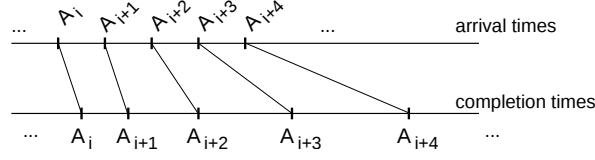


Figure 6: Server Lag Example

We detect server lag by monitoring the average response time and the QoD of the served results. Specifically, we compute the *rate of change* between consecutive calls to the $\text{OVIS}(\theta)$ algorithm. We conclude that server lag is imminent, if:

- the rate of increase for the average response time is too high (for example, a 100 msec increase in average response time over 1000 accesses), or,
- the rate of decrease for the average QoD is too high (for example, a 0.1 drop in QoD over 1000 accesses).

A sudden increase in the average response time is a textbook case for server lag. A sudden decrease in the average QoD indicates that our system, with the current configuration of materialization policies, has surpassed its capacity to handle updates in a timely manner, as a result of server lag.

Server lag is used to detect *infeasible QoD thresholds*. For example a QoD threshold very close to 1 will most likely lead to a server meltdown and should be detected, since no WebView could be materialized and thus the system will be vulnerable to overloads.

OVIS(θ) - QoD Surplus	
0.	$\text{qod_diff} = \text{QoD} - \theta > 0$
1.	ignore all materialized WebViews
2.	ignore all WebViews with $\Delta_{\text{perf}} \geq 0$
3.	find W_i with min Δ_{perf}
4.	if MAX_CHANGE not reached and $(\text{qod_diff} + \Delta_{\text{QoD}}(W_i)) > 0$ then
5.	materialize W_i
6.	$\text{qod_diff} += \Delta_{\text{QoD}}(W_i)$
7.	goto step 3
8.	else
9.	STOP

Figure 7: Pseudo-code for OVIS(θ) - QoD Surplus

Pseudo-code

The $\text{OVIS}(\theta)$ algorithm is in Passive Mode most of the time, collecting statistics (Figure 5). Periodically, $\text{OVIS}(\theta)$ enters Active Mode in order to adapt the materialization plan. Before deciding on a new materi-

OVIS(θ) - QoD Deficit	
0.	$qod_diff = \theta - QoD > 0$
1.	ignore all WebViews not materialized
2.	ignore all WebViews with $\Delta_{QoD} \leq 0$
3.	find W_i with max Δ_{QoD}
4.	if MAX_CHANGE not reached then
5.	stop materializing W_i
6.	$qod_diff -= \Delta_{QoD}(W_i)$
7.	if $qod_diff > 0$
8.	goto step 3
9.	else
10.	STOP
11.	else
12.	STOP

Figure 8: Pseudo-code for OVIS(θ) - QoD Deficit

alization plan, the algorithm will check if there is server lag. If server lag is detected, OVIS(θ) makes all WebViews materialized. This action corresponds to pressing a “panic” button.

Making all WebViews materialized will have the best performance and thus should help alleviate server backlog before it is too late. Materialization essentially “protects” accesses from overload by removing the handling of the updates from the critical path. Assuming “well-behaved” update processes, a surge in updates will lead to reduced QoD without impact on performance.

There are two cases when OVIS(θ) skips an opportunity to adapt the materialization plan:

1. for an initial *warm-up* period we forbid adaptation in order to collect enough statistics about the workload, and
2. after detecting server lag, we impose a short mandatory *cool-down* period, during which we do not allow any plan adaptations, in order to let the system reach a stable state again.

Figures 7 and 8 present the active mode of the OVIS(θ) algorithm under Surplus and Deficit conditions.

Implementing OVIS on a Real System

Although we have not yet implemented OVIS as part of a commercial web server, we believe that doing so would not be very difficult. There are two main issues that need to be addressed: (1) how to invalidate cached WebViews, and (2) how to collect statistics about QoD and response times at the server.

In order to recognize which of the WebViews cached inside the Asynchronous Cache have been invalidated, we propose to rely on the existing replication facilities of modern database management systems [LGZ04]. In order to collect the required statistics, we propose to instrument the Web server appropriately. We have done so successfully for our WebView Materialization study [LR00], by instrumenting the page request module of the Apache Web Server to collect response time information.

4 Experiments

In order to study the online view selection problem, we built *OSim*, a data-intensive web server simulator in C++. The database schema, the costs for updating relations, the costs for accessing/refreshing views, the incoming access stream, the incoming update stream and the level of multitasking are all inputs to the simulator. The simulator processes the incoming access and update streams and generates the stream of responses to the access requests, along with timing information. Among other statistics, the simulator maintains the QoD metric for the served data.

OSim runs in two modes: *static mode* and *adaptive mode*. In static mode, the materialization plan is pre-specified and is fixed for the duration of the simulation. In adaptive mode, the materialization plan is modified at regular intervals using the $OVIS(\theta)$ algorithm (Figures 7 and 8). We report the average response time and the observed QoD for each experiment.

We used synthetic workloads in all experiments. The database contained 200 relations, 500 WebViews and 300 web pages. Each relation was used to create 3-7 WebViews, whereas each web page consisted from 10 to 20 WebViews. Access requests were distributed over web pages following a Zipf-like distribution [BCF⁺99] and the updates were distributed uniformly among relations. We also generated random Web Page Derivation Graphs (like the one in Figure 2). Although updates were distributed uniformly among relations, this did not correspond to uniform distribution of updates to WebViews because of the random view derivation hierarchy. Interarrival rates for the access and the update stream approximated a negative exponential distribution. The cost to update a relation was 150 ms, the cost to access a WebView from the Asynchronous Cache was 10 ms and the cost to generate/refresh a WebView was 150 ms in all experiments.

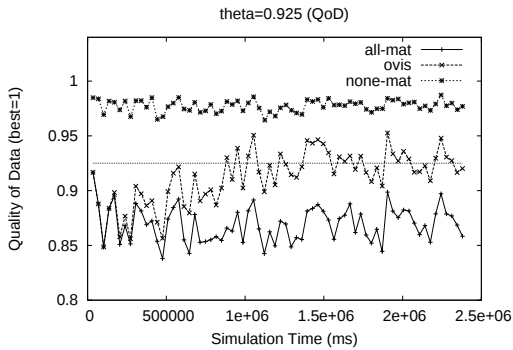


Figure 9: QoD over time for OVIS(0.925)

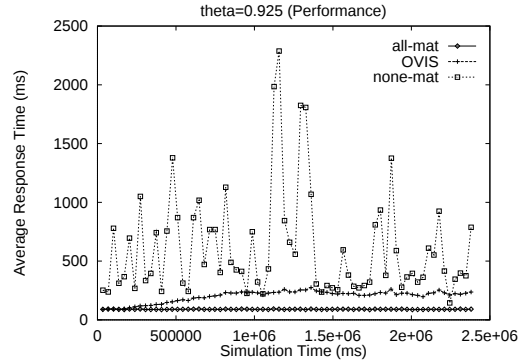


Figure 10: Performance for OVIS(0.925)

4.1 Providing the full spectrum of QoD

In this set of experiments we vary the QoD threshold θ in order to produce the full spectrum of choices between the (low QoD, high performance) case of full materialization and the (high QoD, low performance) case of no materialization. The workload had 35,000 accesses and 32,000 updates. The duration of the experiment was 2400 seconds whereas the QoD threshold θ was set to 0.925.

Figure 9 shows the QoD over time. The top line is the QoD for the no materialization case (i.e. only caching) and the bottom line is the QoD for the fully materialized case. Both policies correspond to static materialization plans. The middle line is the QoD over time for the OVIS algorithm (our adaptive policy) and the straight line is the QoD threshold, 0.925. Initially all WebViews under OVIS start as being materialized. However, in this experiment, the QoD for OVIS quickly “climbs” to the threshold levels and stays around the threshold for the duration of the experiment.

Figure 10 shows the fluctuation of average response time. The top line is the average response time for the no materialization case and the bottom line is the average response time for the fully materialized case. The middle line is the average response time for the OVIS algorithm, with $\theta = 0.925$. The high QoD with the no materialized case is “penalized” by widely fluctuating response times (up to ten times worse than OVIS). On the other hand, the near-constant response times for the fully materialized case correspond to relatively poor QoD. Clearly, the OVIS algorithm provides a good trade-off between these two extremes.

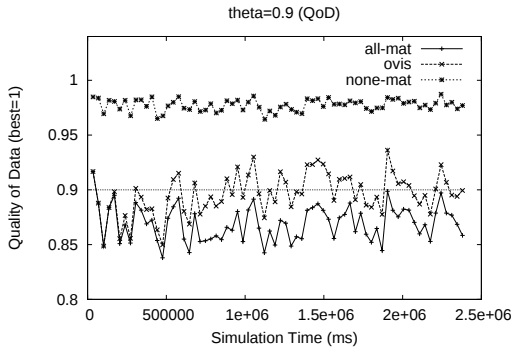


Figure 11: QoD for OVIS(0.90)

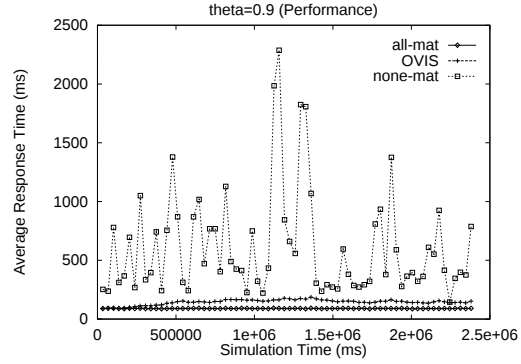


Figure 12: Performance for OVIS(0.90)

We changed the QoD threshold to 0.90 and ran the same experiment. We plot the QoD and the average response times in Figures 11 and 12. In this experiment, the QoD produced by the OVIS(0.90) algorithm is less than that of the OVIS(0.925) algorithm (from Figure 9). In other words, OVIS seems to track the specified QoD threshold.

Finally, we ran the same experiment with a threshold value of 0.85 for OVIS. Since the QoD for the fully-materialized policy is very close to 0.85, the behavior of OVIS mirrored that of the fully-materialized case. This was true for both the data freshness and the average response time.

4.2 Detecting infeasible QoD thresholds

In this set of experiments we wanted to see the behavior of the OVIS algorithm under server lag conditions. The workload had 40,000 accesses and 35,000 updates. The duration of the experiment was 2400 seconds. Under this workload, without materialization, the server exhibits significant lag, the average response times increase monotonically, and the server essentially crashes under the heavy load.

Figure 13 has the average response time for the three policies: no materialization (cached), full materialization (mat) and that produced by the adaptive OVIS algorithm (ovis). The response times for the

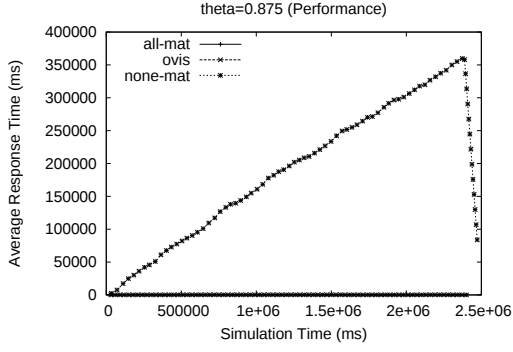


Figure 13: Performance for OVIS(0.875)

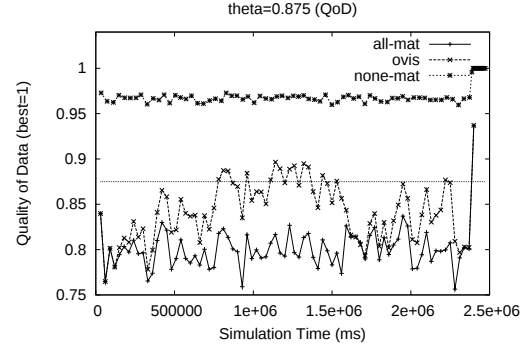


Figure 14: QoD over time for OVIS(0.875)

fully materialized and OVIS are very close together, at the bottom of the graph. The average response time for the no materialization policy increases constantly because the server has been saturated. At the end of the experiment, the average response times without materialization are three orders of magnitude worse than the OVIS or fully-materialized policies. Clearly, this is a situation we want to avoid, regardless of how good the QoD is under the no-materialization policy.

We plot the QoD for the three different policies in Figure 14. The top line is when we do not materialize any WebView, the bottom one is when we materialize all WebViews and the middle line is when we use the OVIS algorithm to decide the materialization policy for each WebView. The OVIS algorithm tries to “climb” towards the QoD threshold 0.875, but, after a while ($t=1563$ sec), detects the server lag and “resets” to a fully-materialized policy, from which it starts to improve the overall QoD again. This behavior is more clear if we look at the average response times of just the fully materialized and the OVIS algorithm, in Figure 15. The response time under the OVIS algorithm increases slowly, then it stabilizes (when the QoD is also stabilized around the QoD threshold), but at some point a sudden increase in response time leads to server lag detection and thus reverting to a fully materialized policy.

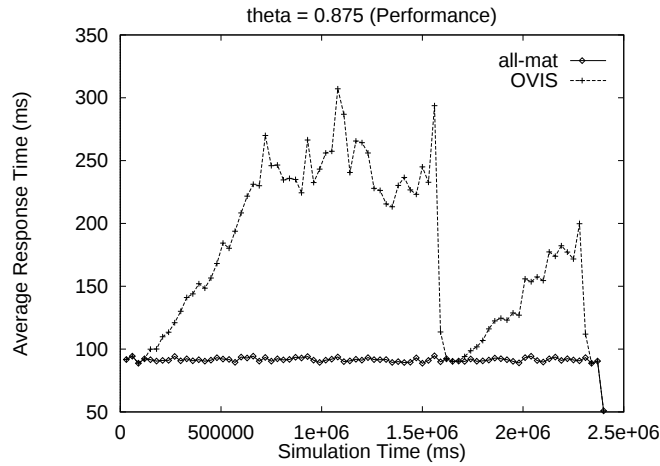


Figure 15: Performance for OVIS(0.875)

We ran the same experiment with different QoD thresholds, one higher (0.90) and one lower (0.85). In the experiment with the lower θ , the QoD stabilizes around the threshold and we do not detect any server lag (Figure 16). On the other hand, in the experiment with the higher θ , the OVIS algorithm detects server lag twice (while trying to reach the high QoD threshold) and resets to a fully materialized policy, at $t=932$ sec and at $t=1834$ sec.

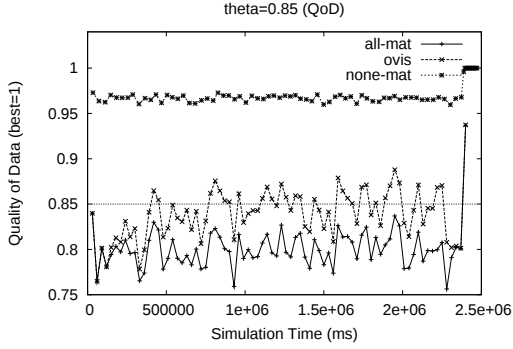


Figure 16: QoD for OVIS(0.85)

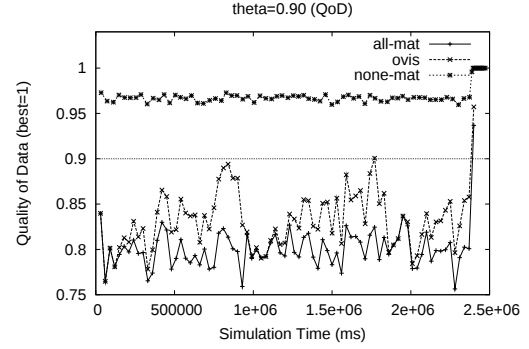


Figure 17: QoD for OVIS(0.90)

4.3 Scaling the number of WebViews

In this set of experiments we evaluated the behavior of the OVIS algorithm with a higher number of WebViews. The workload had 800 relations, 2000 WebViews, and 1500 web pages. Each web page consisted of 10 to 20 WebViews, whereas each relation was used to generate 5 to 15 WebViews. A total of 40,000 web page accesses and 25,000 relation updates occurred in a 2400-second interval.

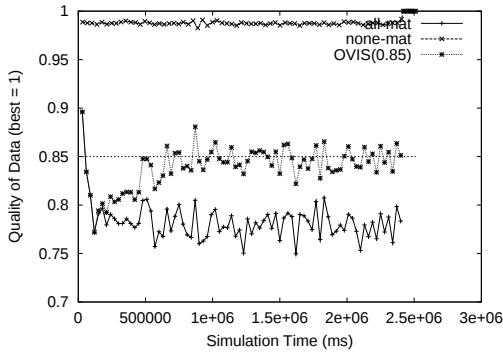


Figure 18: QoD - 25K updates

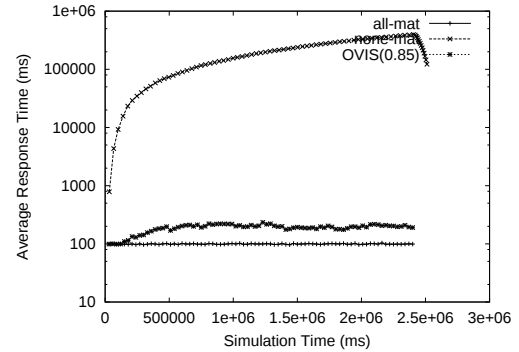


Figure 19: Performance - 25K updates

We plot the QoD and average response time in Figure 18 and Figure 19. In both plots, the top line is the case without any materialization (*none-mat*), the middle curve is when we ran OVIS with a threshold of 0.85 (*ovis(0.85)*), and the bottom line is the fully materialized case (*all-mat*). In Figure 19, the Y-axis (average response time) is in logarithmic scale in order to distinguish between the OVIS and *all-mat* curves. Clearly, even at a higher number of WebViews, OVIS continues to provide a hybrid solution between the

fully-materialized and no materialization cases. Also, OVIS avoids the server overload that is exhibited by the materialize-nothing approach (top curve). In fact, OVIS is able to track the user-specified QoD threshold of 0.85 very well.

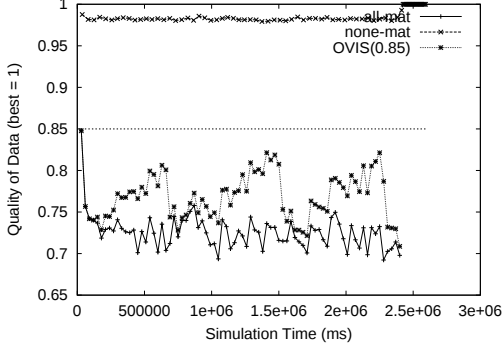


Figure 20: QoD - 30K updates

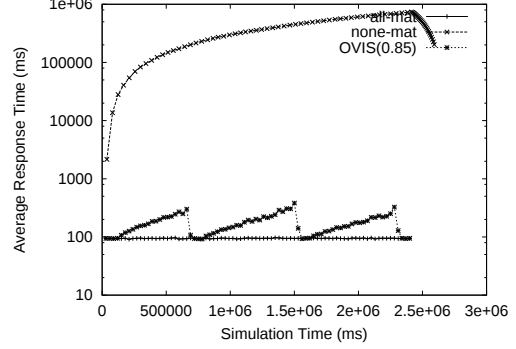


Figure 21: Performance - 30K updates

We ran another set of experiments, with the same setup, but with a heavier update workload: 30,000 relation updates instead of 25,000. We plot the QoD and average response time for this set of experiments in Figure 20 and Figure 21. In both plots, the top line is the case without any materialization (*none-mat*), the middle curve is when we ran OVIS with a threshold of 0.85 (*ovis(0.85)*), and the bottom line is the fully materialized case (*all-mat*). The Y-axis (average response time) is in logarithmic scale for Figure 21 in order to distinguish between the OVIS and *all-mat* curves. In this set of experiments, the user-specified QoD threshold of 0.85 is *infeasible*: although OVIS attempts to reach the threshold, the system detects server lag and reverts to the fully-materialized policy thus avoiding meltdown. The behavior is evident by the zig-zag on the Quality of Data (Figure 20) and the average response time (Figure 21). Nevertheless, OVIS(0.85) still provides a hybrid solution, with higher QoD than the fully-materialized case, and slightly worse average response time. Without materialization we have server overload, resulting in average response times that are three to four orders of magnitude worse than those of the OVIS algorithm.

5 Alternative QoD Metrics

In this section we present an overview of QoD metrics that can be used to measure the freshness of dynamic web pages. We distinguish three dimensions of such metrics: how individual fragment freshness is computed, how it is aggregated at the page level, and how it is aggregated over multiple accesses. We present the different alternatives for each dimension in the following paragraphs.

5.1 Fragment freshness

Measuring the freshness of an individual WebView (or HTML fragment) is the most fundamental component of any QoD metric. We classify such Fragment Freshness Metrics (FFMs) into two categories:

- **boolean**, when a true/false answer to the “is this fragment stale?” question is used. In boolean FFM, the assigned value is typically 0 if the fragment is stale and 1 if it is fresh. Note that there are no intermediate values.

We used a boolean FFM in the presentation of the $\text{OVIS}(\theta)$ algorithm, although the algorithm will work with trivial modifications with any type of FFM, since $\text{OVIS}(\theta)$ is based on QoD differentials and not on the actual values for QoD.

- **numeric**, when a numeric value is used to further specify how much fresh a certain fragment is. In numeric FFM, the assigned value can either be bounded (e.g., when the FFM is in the $[0, 1]$ range), or be completely unbounded. We further explore the different types of numeric FFM in the next paragraphs. We adopt the terminology from [OW02] in our presentation.

Time-based numeric FFM use the time elapsed from the previous update to quantify how stale a certain data item (or fragment in our case) is. Such metrics are especially useful in distributed environments where the exact time of the next update is not known, but instead Time-to-Live (TTL) information is used as an approximation.

One problem with this approach is that the resulting values are unbounded and also their interpretation would vary according to applications and data domains. For example, a stock quote that is 30 minutes old would be practically useless, whereas temperature information that is 30 minutes old is still considered as a reasonable approximation of the current temperature.

We propose to alleviate this problem by introducing a *mapping* of the time-since-last-update to a 0-1 range, which can be defined on a per-application basis. This mapping, which is defined similarly to the QoS curves from Aurora [CCC⁺02], is as follows. There is an initial period after the last update, t_v for which the value is considered valid and the data item is fresh (and thus the FFM has a value of 1). After this period, the “freshness” of the data item declines according to a function (can be linear or any other monotonically decreasing function). At time t_{nv} the freshness of the data item drops to 0. After that time, the freshness remains 0. Figure 22 illustrates the concept. The major advantage of this proposal is that it

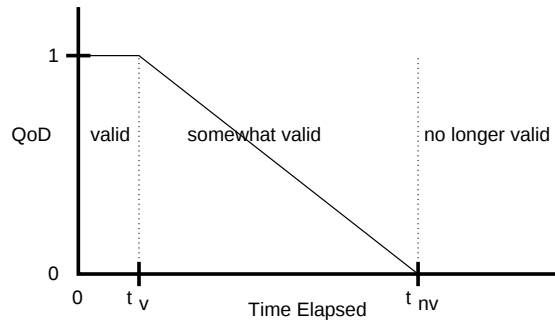


Figure 22: Mapping of time-based FFM to $[0,1]$ space

maps the unbounded time-based metric to an intuitive $[0,1]$ -metric, which can be customized for each specific application domain (by modifying the values of t_v and t_{nv}). The only disadvantage is the small overhead

in computing the metric.

Lag-based numeric FFMs use the number of unapplied updates to quantify how stale a certain fragment is. Such metrics are especially useful when we have exact information on the upcoming updates (i.e. when invalidations are propagate from the origin servers). Similarly to the time-based case, the currently used lag-based FFMs are unbounded and as such are not very intuitive.

In [LFQ04] we proposed measuring the “freshness” of an item as a decreasing function of the number x of updates missed, and specifically:

$$freshness(x) = f(x) = a^x \quad x = 0, 1, \dots \quad (14)$$

We define a as the *freshness decay rate*, which has a value between 1 and 0. The value $a=1$ corresponds to the “do not care” case: the user considers the object perfectly fresh, no matter how many updates it is lagging. The value $a=0$ corresponds to an extremely demanding user, who considers as useless everything, unless it is perfectly up-to-date (we assume that $0^0 = 1$). Notice that each object i has its own freshness decay rate a_i (for example, one object is a stock price that has to be very fresh, while another is a computer manual, which is still useful even if it is slightly out of date).

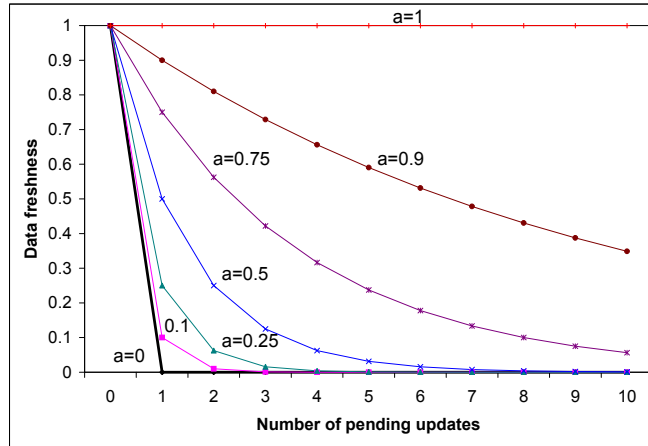


Figure 23: Freshness Decay under various a

Figure 23 illustrates the freshness value under different a values. It is up to the user to set the appropriate value for a and thus determine the behavior of the freshness decay.

The proposed FFM definition has three important properties:

- being a $[0,1]$ metric, it is much more intuitive to use rather than the unbounded divergence/lag metrics which simply count the number of pending updates [YV00, OW02]
- a is a knob that can “record” the varying characteristics of the data items with regards to freshness. It can also be used to record user preferences.

- this model subsumes the typically-used binary model (0 if something is stale, 1 if something is fresh), which has been used by many approaches in the past (e.g. [LR01]). The binary model can be seen as a special case (with $a=0$).

Divergence-based numeric FFMs compare the current version with the most up to date version and quantify the difference in values (i.e. the divergence). Such metrics are especially useful when data items are simple values (e.g. a stock price), but do not work as good on entire HTML fragments, because it is difficult to accurately quantify the difference between two arbitrary HTML/XML fragments.

In the case of simple values one can use the absolute value difference as the divergence metric [SRS02], or normalize the difference (by dividing with the current value) to compute a relative percentage. Obviously, the relative percentage approach cannot be used in cases where the value domain includes 0. Finally, it should be noted that in both the absolute and relative cases the resulting FFM will be unbounded, as even the relative percentage can be more than 100% for large differences.

5.2 Aggregating at the page level

Given a Fragment Freshness Metric, FFM, the question arises how to combine it for fragments that are part of the same web page. There are two main approaches:

- the **minimum** value of the FFM of the component fragments is used when a “guarantee” is needed on the freshness of all fragments on the page. This is especially useful when used in combination with numeric FFM. For boolean metrics, it will lead to an all-or-nothing behavior, which is usually undesirable.
- the **average** value of the FFM of the component fragments is used when a strict guarantee is not needed, but instead we want to get an indication of the freshness of the entire page. This can be used in combination with boolean FFM, to indicate what percentage of fragments should be fresh. This is the approach taken by the OVIS(θ) algorithm.

In the general case, the average can be *weighted*, to indicate the different levels of importance of the fragments composing the web page.

In a system where nesting of fragments is supported, not all fragments should be considered when computing the freshness at the page level. Assume that we have an arbitrary hierarchy of fragments, expressed as a directed acyclic graph. The nodes with zero out-degree are the web pages. Clearly, in such a case, only the freshness of the leaf fragments should be considered when computing the freshness at the page level. If we have a leaf fragment A, then all other fragments that are derived from it will “inherit” the freshness value of A. This is also true in the case of a fragment being composed of multiple lower-level fragments. Therefore, when nesting of fragments is supported only the inner-most fragments must be considered when computing the freshness at the page level.

5.3 Aggregating over multiple access requests

Given an FFM and a way to aggregate it at the page level, the question arises how to combine it for multiple web page accesses. There are two main approaches:

- using the **minimum** value of page freshness will provide a guarantee on the freshness of the data returned back to the user. This is especially useful when combined with the minimum approach of aggregating FFM at the page level.
- using the **average** value of page freshness will not provide a guarantee, but instead provide what is the freshness for the general case and thus allow fluctuations. This is the approach taken by the OVIS(θ) algorithm.

Both approaches can be combined with either approach for aggregating freshness at the page level, with different semantics for each combination.

6 Extensions to the OVIS Algorithm

In this section we outline how the OVIS(θ) algorithm can be extended to handle alternate QoD metrics and QoD requirements from multiple users.

6.1 Supporting alternate QoD metrics

For the presentation of the OVIS(θ) algorithm, we used a boolean Fragment Freshness Metric, FFM, with aggregation using averages both at the page level and among multiple access requests. Since the basic idea behind our algorithm is based on performance/QoD differentials, as such, the algorithm is not tied to any particular FFM. Any boolean or numeric FFM could be used, as long as the differentials of Eq 12 and Eq 13 can be easily computed.

The current implementation of OVIS(θ) assumes that FFMs are aggregated using averages. If we defined QoD to provide strict guarantees, i.e. using minimum instead of average for aggregating at the page level and among multiple access requests, then the OVIS(θ) algorithm must change as follows. Assume that the required QoD threshold is θ_1 . At every adaptation point, we must go through all WebViews and compare each one's QoD against θ_1 . For those below θ_1 , we must act in the same way as in the QoD Deficit case of Figure 8. For those that are above θ_1 , we must act in the same way as in the QoD Surplus case of Figure 7. In other words, we can only “invest” our QoD surplus for those WebViews that are strictly above the threshold, and we must also compensate for any WebViews that are below the threshold. The *MAX_CHANGE* constraint will still be applicable, so we must first deal with the QoD Deficit cases.

6.2 Handling multiple user requirements

In the current version of the OVIS(θ) algorithm, only a single QoD threshold θ was used. We propose to extend OVIS to handle multiple users as follows. Assume that we have n users, with different QoD threshold

requirements $\theta_1, \theta_2, \dots, \theta_n$. The only changes needed to the OVIS algorithm to support multiple users are two:

- Pick the most stringiest of these QoD requirements and use that as the system-wide QoD threshold, i.e. $\theta = \max \theta_i$. This approach is similar to that used for a hierarchy of caches in [SDR03].
- If server lag is detected, then this means that the most stringiest QoD threshold is *infeasible*. In this case, after reverting to the all-materialized state, we must reset the system-wide QoD threshold to be the second stringiest QoD threshold among all users. If another server lag is detected, we repeat the process, further reducing the QoD threshold until a feasible threshold is reached.

7 Related Work

Our work stands between: 1) view selection and run-time buffer management for data warehouses, and 2) dynamic web caching.

View selection has been studied extensively in the context of data warehouses [Rou82, Gup97, GM99, LSTV99, TLS01, MRSR01]. However, in all of the current literature, the selection process is off-line, requiring complete knowledge of the access and update workloads in advance. This is an unrealistic assumption for web servers, which are always online.

Research in run-time buffer management for data warehouses is closer to our work, because the proposed algorithms are online (i.e. they do not require knowledge of the entire access and update stream).

[Sel88] deals with caching the results of frequent or expensive queries in secondary storage. A cost model is presented and ranking-based replacement policy is introduced. Cached results are updated *on demand*, i.e. the next time they are requested in a query.

[SSV96] presents WATCHMAN, which is a data warehouse cache manager. Cache admission and cache replacement are integrated using a “profit metric” that considers the access rate, the size and execution cost of the cached results. This work however does not deal with warehouse updates.

[KR99] presents DynaMat, a data warehouse cache manager that unifies the view selection and view maintenance problems under a single framework. Materialized views are stored in a *view pool* and are refreshed during predetermined update windows. Replacement decisions are made either when the pool becomes full (space-bound case) or because there might not be enough time within the update window to refresh some of the materialized results (time-bound case). Updates in DynaMat are *off-line*, since queries are not answered during the update window.

Finally, [RBSS02] present a scheduling framework that takes into account freshness when it decides to replicate OLAP data or route transactions in a database cluster.

Performing updates concurrently with user queries is the main difference between existing work on materialized view selection and our work. Current state of the art in materialized view selection aims to improve overall performance (QoS). However, in the web server environment, updates are online and thus we need to consider both QoS and Quality of Data (QoD).

Dynamic web caching was introduced in [IC97]. Until recently, research has focused on providing an infrastructure to support caching of dynamically generated web pages [CIW⁺00, YFIV00]. The decision of which pages to cache, when to cache them and when to invalidate or refresh them is left to the application program or the web site designer. There is recent work on cache management for dynamic web content. [DDT⁺01, DDT⁺02] presents a *Dynamic Content Accelerator* prototype that can cache fragments of dynamically generated web pages. [LKM⁺02, ABK⁺03] present DBCache, which is an IBM DB2 database which can cache entire database tables transparently to the application server.

The biggest problem of employing caching techniques for dynamic web content is having the handling of updates in the critical path of serving access requests which can have detrimental effects on server performance. In contrast, OVIS(θ) will effectively mix caching with view materialization to strike the best balance between performance and data freshness, while avoiding server backlogs.

8 Conclusions

Traditional caching techniques, if used in isolation to accelerate dynamic web content, face the possibility of server backlogs, because the handling of updates is in the critical path of serving access requests. In this paper, we have introduced the Online View Selection Problem: dynamically select which views to materialize in order to maximize performance while keeping data freshness at acceptable levels. We presented OVIS(θ), an adaptive algorithm which combines view materialization with caching, and effectively allows for the decoupling of serving of access requests and handling of updates. Parameter θ in OVIS is the level of data freshness that is considered acceptable for the current application. Through extensive experiments we showed that OVIS(θ) can: 1) provide the full spectrum of quality of data, and 2) detect and prevent server backlogs. We envision OVIS(θ) being used together with current dynamic content accelerators in order to build web-aware database servers that are self-manageable, robust, and scalable.

References

- [ABK⁺03] Mehmet Altinel, Christof Bornhovd, Sailesh Krishnamurthy, C. Mohan, Hamid Pirahesh, and Berthold Reinwald. “Cache Tables: Paving the Way for an Adaptive Database Cache”. In *Proceedings of the 29th Very Large Data Bases (VLDB) Conference*, pages 718–729, Berlin, Germany, September 2003.
- [BAK⁺03] Christof Bornovd, Mehmet Altinel, Sailesh Krishnamurthy, C. Mohan, Hamid Pirahesh, and Berthold Reinwald. “DBCache: Middle-tier Database Caching for Highly Scalable e-Business Architectures”. In *Proc. of ACM SIGMOD*, page 662, San Diego, California, USA, June 2003.
- [BCF⁺99] Lee Breslau, Pei Cao, Li Fan, Graham Phillips, and Scott Shenker. “Web Caching and Zipf-like Distributions: Evidence and Implications”. In *Proceedings of 1999 INFOCOM Conference*, pages 126–134, New York, USA, March 1999.

- [CCC⁺02] Don Carney, Ugur Cetintemel, Mitch Cherniack, Christian Convey, Sangdon Lee, Greg Seidman, Michael Stonebraker, Nesime Tatbul, and Stan Zdonik. “Monitoring Streams: A New Class of Data Management Applications”. In *Proceedings of the 28th Very Large Data Bases (VLDB) Conference*, pages 215–226, Hong Kong, China, August 2002.
- [CIW⁺00] Jim Challenger, Arun Iyengar, Karen Witting, Cameron Ferstat, and Paul Reed. “A Publishing System for Efficiently Creating Dynamic Web Content”. In *Proceedings of 2000 INFOCOM Conference*, pages 844–853, Tel Aviv, Israel, March 2000.
- [DDT⁺01] Anindya Datta, Kaushik Dutta, Helen M. Thomas, Debra E. VanderMeer, Krithi Ramamritham, and Dan Fishman. “A Comparative Study of Alternative Middle Tier Caching Solutions to Support Dynamic Web Content Acceleration”. In *Proceedings of the 27th Very Large Data Bases (VLDB) Conference*, pages 667–670, Rome, Italy, September 2001.
- [DDT⁺02] Anindya Datta, Kaushik Dutta, Helen M. Thomas, Debra E. VanderMeer, Suresha, and Krithi Ramamritham. “Proxy-Based Acceleration of Dynamically Generated Content on the World Wide Web: An Approach and Implementation”. In *Proc. of ACM SIGMOD*, pages 97–108, Madison, Wisconsin, USA, June 2002.
- [GM99] Ashish Gupta and Inderpal Singh Mumick, editors. “*Materialized Views: Techniques, Implementations, and Applications*”. MIT Press, June 1999.
- [Gup97] Himanshu Gupta. “Selection of Views to Materialize in a Data Warehouse”. In *Proceedings of ICDT*, pages 98–112, Delphi, Greece, January 1997.
- [IC97] Arun Iyengar and Jim Challenger. “Improving Web Server Performance by Caching Dynamic Data”. In *Proceedings of the USENIX Symposium on Internet Technologies and Systems*, Monterey, CA, USA, December 1997.
- [Jac88] V. Jacobson. “Congestion Avoidance and Control”. In *Proceedings of ACM SIGCOMM*, pages 314–329, Stanford, California, USA, August 1988.
- [Jai91] Raj Jain. “*The Art of Computer Systems Performance Analysis*”. John Wiley & Sons, 1991.
- [KR99] Yannis Kotidis and Nick Roussopoulos. “DynaMat: A Dynamic View Management System for Data Warehouses”. In *Proc. of ACM SIGMOD*, pages 371–382, Philadelphia, USA, June 1999.
- [LFQ04] Alexandros Labrinidis, Christos Faloutsos, and Huiming Qu. “MAXF: A QoD-Aware Cache Refresh Policy Under Online Updates”. submitted for publication, 2004.
- [LGZ03] Paul Larson, Jonathan Goldstein, and Jingren Zhou. “Transparent Mid-Tier Database Caching in SQL Server”. In *Proc. of ACM SIGMOD*, page 661, San Diego, California, USA, June 2003.

- [LGZ04] Paul Larson, Jonathan Goldstein, and Jingren Zhou. “Transparent Mid-Tier Database Caching in SQL Server”. In *Proceedings of the Twentieth International Conference on Data Engineering*, pages 177–189, Boston, MA, USA, April 2004.
- [LKM⁺02] Qiong Luo, Sailesh Krishnamurthy, C. Mohan, Hamid Pirahesh, Honguk Woo, Bruce G. Lindsay, and Jeffrey F. Naughton. “Middle-tier Database Caching for e-Business”. In *Proc. of ACM SIGMOD*, page 662, Madison, Wisconsin, USA, June 2002.
- [LR00] Alexandros Labrinidis and Nick Roussopoulos. “WebView Materialization”. In *Proc. of ACM SIGMOD*, pages 367–378, Dallas, Texas, USA, May 2000.
- [LR01] Alexandros Labrinidis and Nick Roussopoulos. “Update Propagation Strategies for Improving the Quality of Data on the Web”. In *Proceedings of the 27th Very Large Data Bases (VLDB) Conference*, pages 391–400, Rome, Italy, September 2001.
- [LSTV99] Spyros Ligoudistianos, Timos K. Sellis, Dimitri Theodoratos, and Yannis Vassiliou. “Heuristic Algorithms for Designing a Data Warehouse with SPJ Views”. In *Proceedings of the First International Conference on Data Warehousing and Knowledge Discovery (DaWaK)*, pages 96–105, Florence, Italy, August 1999.
- [MRSR01] Hoshi Mistry, Prasan Roy, S. Sudarshan, and Krithi Ramamritham. “Materialized View Selection and Maintenance Using Multi-Query Optimization”. In *Proc. of ACM SIGMOD*, pages 307–318, Santa Barbara, California, USA, May 2001.
- [OW02] Chris Olston and Jennifer Widom. “Best-Effort Cache Synchronization with Source Cooperation”. In *Proc. of ACM SIGMOD*, pages 73–84, Madison, Wisconsin, USA, June 2002.
- [RBSS02] U. Rohm, K. Bohm, H.-J. Schek, and H. Schuldt. “FAS - A Freshness-Sensitive Coordination Middleware for a Cluster of OLAP Components”. In *Proceedings of the 28th Very Large Data Bases (VLDB) Conference*, Hong Kong, China, August 2002.
- [Rou82] Nick Roussopoulos. “View Indexing in Relational Databases”. *ACM Transactions on Database Systems*, 7(2):258–290, June 1982.
- [SDR03] Shetal Shah, Shyamshankar Dharmarajan, and Krithi Ramamritham. “An Efficient and Resilient Approach to Filtering and Disseminating Streaming Data”. In *Proceedings of the 29th Very Large Data Bases (VLDB) Conference*, pages 57–68, Berlin, Germany, September 2003.
- [Sel88] Timos Sellis. “Intelligent caching and indexing techniques for relational database systems”. *Information Systems*, 13(2):175–185, 1988.
- [SRS02] Shetal Shah, Krithi Ramamritham, and Prashant J. Shenoy. “Maintaining Coherency of Dynamic Data in Cooperating Repositories”. In *Proceedings of the 28th Very Large Data Bases (VLDB) Conference*, pages 526–537, Hong Kong, China, August 2002.

- [SSV96] Peter Scheuermann, Junho Shim, and Radek Vingralek. “WATCHMAN : A Data Warehouse Intelligent Cache Manager”. In *Proc. of VLDB Conference*, pages 51–62, Bombay, India, September 1996.
- [TLS01] Dimitri Theodoratos, Spyros Ligoudistianos, and Timos K. Sellis. “View selection for designing the global data warehouse”. *Data and Knowledge Engineering*, 39(3):219–240, 2001.
- [YFIV00] Khaled Yagoub, Daniela Florescu, Valerie Issarny, and Patrick Valduriez. “Caching Strategies for Data-Intensive Web Sites”. In *Proceedings of the 26th Very Large Data Bases (VLDB) Conference*, pages 188–199, Cairo, Egypt, September 2000.
- [YV00] Haifeng Yu and Amin Vahdat. “Design and Evaluation of a Continuous Consistency Model for Replicated Services”. In *Proc. of the Fourth Symposium on Operating Systems Design and Implementation*, pages 305–318, October 2000.