

УДК 004.4

СРАВНЕНИЕ ПОСЛЕДОВАТЕЛЬНОЙ И ПАРАЛЛЕЛЬНОЙ ПРОГРАММНЫХ РЕАЛИЗАЦИЙ ГИБРИДНОЙ ЖИДКОСТНОЙ МОДЕЛИ ИНФОРМАЦИОННЫХ ПОТОКОВ В ВЫСОКОСКОРОСТНЫХ МАГИСТРАЛЬНЫХ ИНТЕРНЕТ-КАНАЛАХ

Басавин Д.А., Поршневу С.В.

ФГАОУ ВПО «Уральский федеральный университет имени первого

Президента России Б.Н. Ельцина», Екатеринбург,

e-mail: basavind@gmail.com, e-mail: sergey_porshnev@mail.ru

Описаны результаты сравнения последовательной (на центральном процессоре) и параллельной (на графическом процессоре) программных реализаций гибридной жидкостной модели информационных потоков в магистральном интернет-канале. Данные реализации адаптированы для моделирования информационного потока на входе и выходе отдельного маршрутизатора, нагруженного информационными потоками, обусловленными активностью пользователей. Показано, что количественные характеристики информационных потоков на входе и выходе маршрутизатора (средней скорости передачи данных, максимальной скорости передачи данных, минимальной скорости передачи данных), рассчитанные с помощью каждой из использованных программных реализаций, согласуются друг с другом. Получены зависимости времени вычислений от числа пользователей, свидетельствующие о целесообразности использования параллельной реализации гибридной жидкостной модели при числе пользователей 2^{12} и более.

Ключевые слова: гибридная жидкостная модель, параллельные вычисления, графические процессоры, графические процессоры общего назначения, технология OpenCL

COMPARISON OF SEQUENTIAL AND PARALLEL SOFTWARE IMPLEMENTATIONS OF HYBRID FLUID MODEL OF INFORMATION FLOWS IN BACKBONE NETWORKS

Basavin D.A., Porshnev S.V.

¹Federal State Autonomous Educational Institution

of Higher Professional Education «Ural Federal University

named after the first President of Russia B.N. Yeltsin»,

Ekaterinburg, e-mail: basavind@gmail.com, e-mail: sergey_porshnev@mail.ru

This paper presents results of comparing sequential (on CPU) and parallel (on GPU) software implementations of hybrid fluid model of information flows in backbone networks (such as Internet). These implementations are adapted for information flows modeling at the input and output of single router loaded by information streams due to users activities. It is shown, that quantitative characteristics of information flows at the input and output of router (average data rate, maximum data rate, minimum data rate), calculated by each of the used implementations consistent with each other. The dependences of the computation time due to the number of users showing the feasibility of using the parallel implementation of a hybrid fluid model when the number of users equal and more than 2^{12} .

Keyword: hybrid fluid model, parallel computing, GPU, GPGPU, OpenCL

В современных условиях, когда осуществлен переход на так называемые «безлимитные» тарифные планы, в которых одним из основных критериев качества обслуживания является предоставление заявленной пользователю скорости передачи данных, резко возрастают требования к техническим решениям, принимаемым при проектировании современных телекоммуникационных сетей. При этом одним из основных оказывается вопрос о количестве пользователей, которое необходимо обслуживать при условии обеспечения заявленной скорости доступа к интернет-каналам на имеющемся в наличии сетевом оборудо-

вании. Как показывает практика, сегодня по-прежнему наиболее популярными остаются эмпирические методы оценки технических характеристик сетевого оборудования (например, метод экспертных оценок), которые в ряде случаев оказываются недостаточно точными. Это, в свою очередь, может приводить как к неоправданным затратам, так и к несоответствию между заявленным и фактическим качеством предоставляемых пользователю сетевых сервисов и, как следствие, невыполнению провайдером взятых на себя обязательств.

Анализ известных решений данной проблемы показывает, что подход, основанный

на использовании математических моделей информационных потоков в каналах передачи данных позволяет, с одной стороны, сделать адекватный выбор необходимого сетевого оборудования и при этом не требующий его приобретения, монтажа, настройки и проверочного тестирования на действующих интернет-каналах, с другой. Среди подобных математических моделей можно выделить следующие классы, которым, однако, присущи известные недостатки [6]:

1. Аналитические модели (в первую очередь модели теории массового обслуживания), в которых не удается учесть особенности современных механизмов регулирования скорости (механизмы реакции на потери пакетов, ограничения потоков отдельных пользователей, влияние потоков друг на друга и так далее) [5, 9].

2. Программы-генераторы сетевого трафика (в том числе статистические модели трафика), в которых оказывается невозможным учесть особенности передачи генерируемого трафика по каналу передачи данных, а также учесть механизмы обратной связи при потере пакетов, используемые в современных интернет-каналах для управления скоростью передаваемого потока данных [3, 4, 8, 14].

3. Сетевые пакетные симуляторы – специализированные программные продукты, предназначенные для моделирования каналов с умеренной пропускной способностью (потоки порядка десятков Мбит/с). Данные модели позволяют описать процесс передачи данных по сети (на уровне отдельных пакетов) и учесть механизмы регулирования скорости потоков трафика [11].

4. Жидкостные модели, учитывающие механизмы управления скоростью потоков передачи, что позволяет существенно уменьшить число рассматриваемых событий при моделировании интернет-трафика за счет перехода от рассмотрения процессов распространения в канале передачи данных отдельных пакетов к рассмотрению укрупненных групп пакетов.

Напомним, что в жидкостной модели сетевой трафик описывается в терминах изменения во времени скорости передачи данных i -го потока $W_i(t)$ и длины очереди на входе в l -й канал $q_l(t)$. Данная модель, позволяющая учесть потерю пакетов при их передаче и реализовывать современные алгоритмы управления скоростью потоков, описывается следующей системой дифференциальных уравнений:

$$\frac{dW_i(t)}{dt} = \frac{l_{(W_i(t)-M_i)}}{R_i(t)} - \frac{W_i(t)}{2} \cdot \lambda_i(t);$$

$$\frac{dq_l(t)}{dt} = -l_{q(t)} \cdot C_l + \sum_{i \in N_l} A_i^l(t),$$

где $l_{f(t)}$ – функция Хевисайда:

$$l_{f(t)} = \begin{cases} 1, & \text{если } f(t) \geq 0, \\ 0, & \text{если } f(t) < 0, \end{cases}$$

$R_i(t)$ – время двойного оборота (RTT) i -го потока; $\lambda_i(t)$ – скорость потери пакетов i -го потока; C_l – пропускная способность l -го канала, который обслуживается данным маршрутизатором; $A_i^l(t)$ – скорость переда-

чи i -го потока по l -му каналу, $A_i^l(t) = \frac{W_i^l(t)}{R_i^l(t)}$.

К недостаткам классической жидкостной модели [13] и ее последующим известным модификациям (см., например, [12], авторы которой предложили способ интеграции жидкостной модели трафика и сетевого симулятора ns-2), следует отнести нерешенность проблемы учета рассогласованного (дискретного) характера действий пользователя. Отмеченный недостаток в известной мере был преодолен в [6]. Здесь была предложена гибридная жидкостная модель информационных потоков в магистральном интернет-канале, представляющая собой комбинацию жидкостной модели [13] и статистического варианта модели абстрактных источников трафика. Гибридная модель позволяет описывать трафик в мультисервисных сетях с учетом присущего протоколу TCP механизма обратного влияния загрузки сети на работу источника, а также учитывать современные политики управления скоростью доступа к сети интернет отдельных пользователей.

Данная модель доведена авторами до законченной последовательной (работающей на центральном процессоре – CPU) программной реализации [7] и ее работоспособность подтверждена сходством свойств реального и синтетического интернет-трафиков [6]. В то же время опыт практического использования последовательной реализации гибридной жидкостной модели показал, что временные затраты, необходимые для проведения расчетов, даже варианта сети с одним маршрутизатором и относительно небольшого количества пользователей, оказываются весьма значительными и возрастают вместе с количеством пользователей.

В связи с тем, что один из наиболее эффективных на сегодняшний день подходов, позволяющих повысить скорость вычислений при моделировании компьютерных сетей, состоит в использовании параллельных вычислений, реализованных на

графических процессорах [15], авторы обосновали возможность распараллеливания как «классической», так и гибридной жидкостной моделей [1], и создали параллельную (работающую на графическом

процессоре – GPU) программную реализацию варианта с одним маршрутизатором [2], блок-схема которой представлена на рис. 1. При этом, как очевидно, требуется подтверждение ее работоспособности.

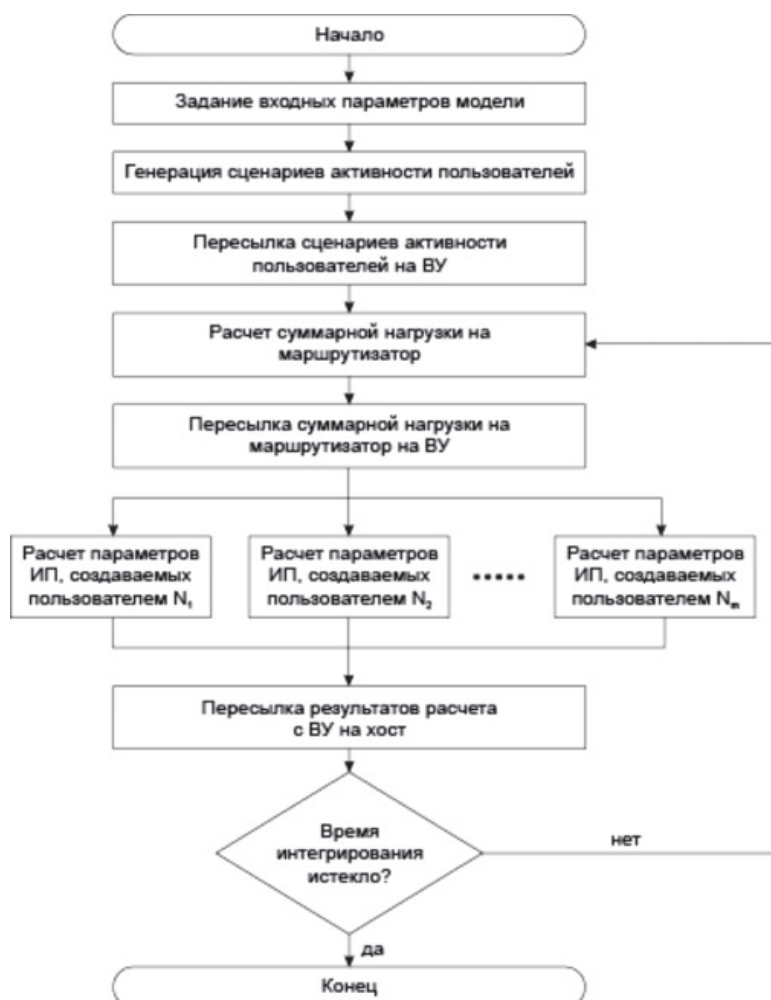


Рис. 1. Блок-схема параллельной жидкостной модели (ВУ – Вычислительное устройство, ИП – информационный поток)

В статье приводятся сравнительный анализ свойств синтетического интернет-трафика, генерируемого последовательной и параллельной программными реализациями гибридной жидкостной модели, а также оценка ускорения расчетов относительно количества пользователей.

Результаты моделирования информационных потоков с помощью параллельной программной реализации гибридной жидкостной модели

В ходе исследований были проведены компьютерные эксперименты, в которых моделировался информационный поток на входе магистрального маршрутизатора,

на вход которого поступают информационные потоки, обусловленные запросами пользователей. Был проведен ряд экспериментов, нацеленных на оценку нагрузки на маршрутизатор, в зависимости от времени между запросами пользователей и их гарантированной скорости доступа, а также в зависимости от количественного соотношения пользователей двух типов («слоны» и «мулы») и пропускной способности маршрутизатора. Описание характеристик экспериментов и их результаты:

1. Зависимости «мгновенных» значений скорости передачи данных в канале при обслуживании $N = 1000$ пользователей в режиме ограничения скорости $CIR = 3$ Мбит/с, при среднем времени между запросами

0,05 и 5 с; а также аналогичная зависимость для случая $CIR = 10$ Мбит/с, при среднем времени между запросами – 5 с (рис. 2). При этом использовались следующие фиксированные значения параметров гибридной модели:

– время моделирования – 600 с;

– количество пользователей – 1000;
– среднее время задержки пользователей (RTT) – от 20 до 70 мс;
– согласованный (BC) и расширенный (BE) всплеск – 50 Кб;
– длительность активности пользователей (период ON) – от 0,3 до 1 с.

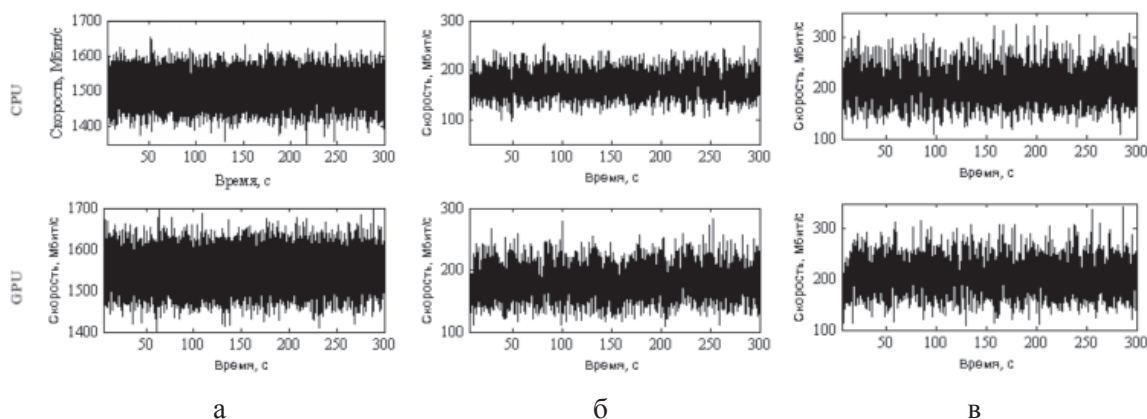


Рис. 2. Зависимости «мгновенных» значений скорости передачи данных в канале при обслуживании $N = 1000$ пользователей в режиме ограничения скорости:
а – $CIR = 3$ Мбит/с, среднее время между запросами = 0,05 с; б – $CIR = 3$ Мбит/с, среднее время между запросами = 5 с; в – $CIR = 10$ Мбит/с, среднее время между запросами = 5 с

2. Зависимости «мгновенных» значений объемов данных, поступающих на вход маршрутизатора от времени при различных сценариях активности пользователей (рис. 3). При этом в качестве источников трафика были использованы следующие классы активности пользователей: «слоны» (пользователи, осуществляющие скачивание и передачу больших объемов данных (сотни Мб), сохраняющие активность на протяжении всего процесса моделирования); «мулы» (пользователи, осуществляющие скачивание и передачу объемов данных

среднего размера (сотни и тысячи Кб), при моделировании длительность периодов активности «мулов» (ON) варьировалась в диапазоне от 5 до 60 с, среднее время между запросами «мулов» – 3 с (период OFF)). Отметим, что класс «мыши», к которому относятся пользователи, осуществляющие скачивание и передачу небольших объемов данных (десятки Кб), в модели не учитывался, поскольку доля трафика, приходящегося на данный класс пользователей, не превосходит 5–10% общего объема трафика в интернет-каналах.

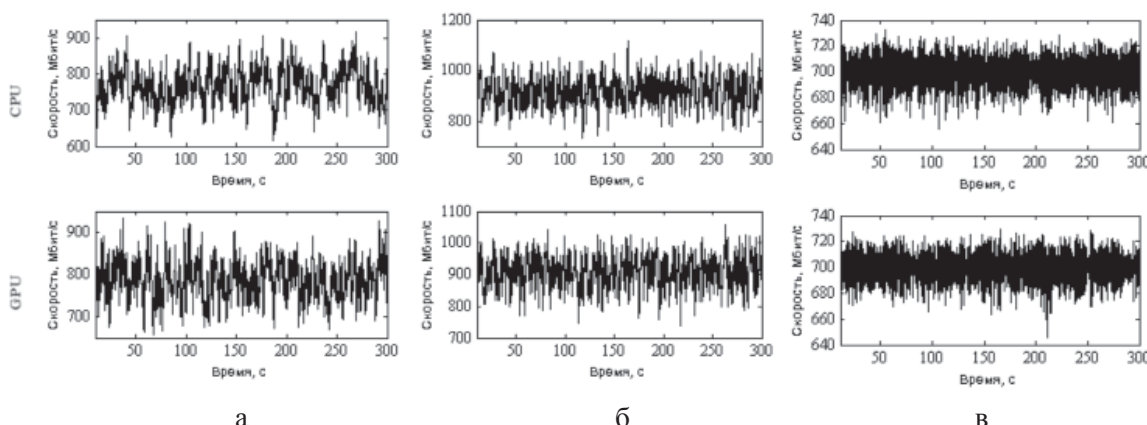


Рис. 3. Зависимости «мгновенных» значений скорости передачи данных в канале при обслуживании $N = 100$ пользователей в режиме ограничения скорости:
а – 10 «слонов», 90 «мулов», пропускная способность маршрутизатора не ограничена;
б – 90 «слонов», 10 «мулов», пропускная способность маршрутизатора не ограничена;
в – 10 «слонов», 90 «мулов», пропускная способность маршрутизатора ограничена 700 Мбит/с

Расчеты проводились при следующих фиксированных значениях параметров гибридной модели:

- время моделирования – 300 с;
- среднее время задержки пользователей (RTT) – от 20 до 100 мс;
- гарантированная скорость доступа (CIR) – 10 Мбит/с;

- согласованный всплеск (BC) – 1000 Кб;
- расширенный всплеск (BE) – 2000 Кб;
- размер буфера сетевого маршрутизатора – 1000 Кб.

Варьируемые параметры моделей представлены в таблице.

К сравнению скорости передачи данных

Эксперимент	Минимальная скорость	Средняя скорость	Максимальная скорость
2а) CPU	1414,9 Мбит/с	1599,8 Мбит/с	1709,5 Мбит/с
2а) GPU	1414,4 Мбит/с	1564,1 Мбит/с	1710,8 Мбит/с
2б) CPU	98,9 Мбит/с	178,4 Мбит/с	269,2 Мбит/с
2б) GPU	100,5 Мбит/с	177,5 Мбит/с	273,1 Мбит/с
2в) CPU	111,7 Мбит/с	209,1 Мбит/с	337,1 Мбит/с
2в) GPU	107,4 Мбит/с	208,8 Мбит/с	333,9 Мбит/с

Сравнительный анализ результатов моделирования, полученных с помощью последовательной и параллельной реализаций гибридной жидкостной модели

Для сравнения результатов моделирования были использованы следующие количественные характеристики: минимальная скорость передачи данных, средняя скорость передачи данных, максимальная скорость передачи данных. Выбранные характеристики для каждой из программных реализаций представляют собой среднее значение из результатов 50 экспериментов и представлены в таблице.

Из таблицы видно, что выбранные параметры согласуются друг с другом, что позволяет сделать вывод о работоспособности разработанной параллельной программной

реализации гибридной жидкостной модели. Незначительные отличия объясняются случайными характеристиками пользователей (время активности и простоя «мулов» и время двойного оборота сигнала), а также погрешностями, вносимыми графическими процессорами, связанными с их упрощенной архитектурой.

Сравнительный анализ скорости расчетов однопроцессорной и параллельной реализации гибридной жидкостной модели

В ходе проведенных исследований также было проведено измерение времени вычислений, затрачиваемых последовательной и параллельной реализацией гибридной жидкостной модели при различном числе пользователей. Зависимости данных показателей от времени представлены на рис. 4.

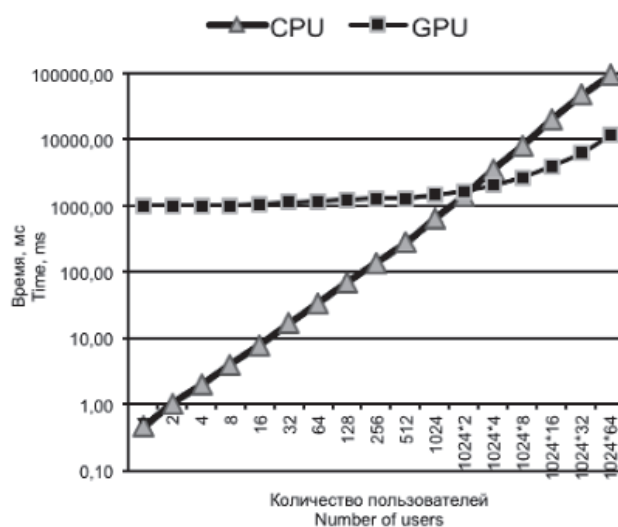


Рис. 4. Зависимость времени моделирования на центральном (CPU) и графическом (GPU) процессорах от количества пользователей

Из рис. 4 видно, что скорость вычислений параллельной реализации гибридной жидкостной модели при большом количестве пользователей ($> 2^{12}$) оказывается выше, чем соответствующий однопроцессорный вариант обсуждаемой модели (например, когда число пользователей составляет 2^{16} – скорость параллельной реализации гибридной жидкостной модели оказывается в 8 раз быстрее последовательной реализации).

В то же время, если число пользователей относительно невелико ($< 2^{10}$), то, напротив, скорость вычислений последовательной реализации гибридной жидкостной модели оказывается выше. Например, при одном пользователе время расчета модели составляет приблизительно 1 секунду.

Данный результат объясняется тем, что вне зависимости от количества пользователей в параллельной реализации гибридной жидкостной модели изначально должны быть инициализированы вычислительные ядра графического процессора, поэтому время, затрачиваемое на данную процедуру, определяет минимально возможное время, за которое модель может быть полностью просчитана. Практически линейная зависимость времени моделирования от числа пользователей при небольшом числе пользователей свидетельствует о неполной загрузке графического процессора, большая часть вычислительных элементов которого простаивает.

Заключение

Проведено сравнение последовательного и параллельного вариантов программных реализаций гибридной жидкостной модели, адаптированных для моделирования информационного потока на входе и выходе отдельного маршрутизатора, нагруженного информационными потоками, обусловленными активностью пользователей.

Сравнение количественных характеристик информационных потоков на входе маршрутизатора (минимальной скорости передачи данных, средней скорости передачи данных, максимальной скорости передачи данных) позволяет сделать вывод о согласованности результатов, получаемых использованными программными реализациями.

Проведено исследование зависимости скорости выполнения вычислений каждой из программных реализаций, результаты которого показывают, что скорость вычислений параллельной реализации гибридной жидкостной модели при большом количестве пользователей ($> 2^{12}$) оказывается выше, чем соответствующий последовательный вариант обсуждаемой модели.

При малом числе пользователей, наоборот, скорость работы однопроцессорного варианта программной реализации гибридной жидкостной модели оказывается выше, что объясняется тем, что независимо от числа пользователей требуется время на инициализацию вычислительных ядер на графическом процессоре.

Список литературы

1. Басавин Д.А., Поршнева С.В. Параллельная гибридная жидкостная модель высокоскоростных информационных потоков в магистральных интернет-каналах // Естественные и технические науки. – 2013. – № 1. – С. 317–326.
2. Басавин Д.А., Поршнева С.В. Свидетельство о государственной регистрации программ для ЭВМ № 2013614980 (Заявка № 2013612648. Дата поступления 01 апреля 2013 г. Дата регистрации в Реестре программ для ЭВМ 27 мая 2013 г.).
3. Генератор трафика ITG Режим доступа. – URL: <http://www.grid.unina.it/software/ITG/link.php> (дата доступа 18.07.2013).
4. Генератор трафика TG. Режим доступа. – URL: <http://www.postel.org/tg/tg.html> (дата доступа 18.07.2013).
5. Гнеденко, Б.В. Введение в теорию массового обслуживания / Б.В. Гнеденко, И.Н. Коваленко. – М.: КомКнига, 2005. – 400 с.
6. Гребенкин М.К., Поршнева С.В. Гибридная жидкостная модель магистрального Интернет-канала // Saarbrücken: LAP LAMBERT Academic Publishing GmbH & Co. KG. – 2012. – 163 с.
7. Гребенкин М.К., Поршнева С.В. Программная реализация гибридной жидкостной модели информационных потоков в высокоскоростных магистральных интернет-каналах // Свидетельство о государственной регистрации программ для ЭВМ № 2012616118 (Заявка № 2012613786. Дата поступления 11 мая 2012 г. Дата регистрации в Реестре программ для ЭВМ 4 июля 2012 г.).
8. Сетевой симулятор OPNET. – Режим доступа: URL: <http://opnet.org> (дата доступа 18.07.2013).
9. Хинчин, А.Я. Математические методы теории массового обслуживания // Тр. Мат. Института им. В.А. Стеклова АН СССР. – 1955.
10. D-ITG, Distributed Internet Traffic Generator. – Режим доступа: URL: <http://www.grid.unina.it/software/ITG/download.php> (дата доступа 18.07.2013).
11. Marsan M.A., Garetto M., Giaccone P., Leonardi E., Schiattarella E., Tarello A. Using Partial Differential Equations to Model TCP Mice and Elephants in Large IP Networks. *INFOCOM 2004. Twenty-third Annual Joint Conference of the IEEE Computer and Communications Societies*. 2004. Vol. 4. P 2821–2832.
12. Misra V., Gong W.-B., Towsley D. Fluid-Based Analysis of a Network of AQM Routers Supporting TCP Flows with an Application to RED // *ACM SIGCOMM Computer Communication Review*. – 2000. – Vol. 30. – Issue 4. – P. 151–160.
13. Tmix: A Tool for Generating Realistic TCP Application Workloads in ns-2. – Режим доступа: URL: <http://www.sigcomm.org/ccr/drupal/files/p67-v36n31-weigle.pdf> (дата доступа 18.07.2013).
14. Xu Zh., Bagrodia R. GPU-Accelerated Evaluation Platform for High Fidelity Network Modeling, *Proceeding PADS '07 Proceedings of the 21st International Workshop on Principles of Advanced and Distributed Simulation*. – 2007. – P. 131–140.

References

1. Basavin D.A., Porshnev S.V. Parallel'naja gibridnaja zhidkostnaja model' vysokoskorostnyh informacionnyh potokov v magistral'nyh internet-kanalah // *Estestvennye i tehnicheckie nauki*, 2013. no. 1. pp. 317–326.

2. Basavin D.A., Porshnev S.V. Svidetel'stvo o gosudarstvennoj registracii programm dlja JeVM no. 2013614980 (Zajavka no. 2013612648. Data postuplenija 01 aprelja 2013 g. Data registracii v Reestre programm dlja JeVM 27 maja 2013 g.)
3. Generator trafika ITG, Available at: <http://www.grid.unina.it/software/ITG/link.php> (accessed at 18.07.2013).
4. Generator trafika TG, Available at: <http://www.postel.org/tg/tg.html> (accessed at 18.07.2013).
5. Gnedenko, B.V. Vvedenie v teoriju massovogo obsluzhivaniya / B.V. Gnedenko, I.N. Kovalenko // M.: KomKniga, 2005. 400 p.
6. Grebenkin M.K., Porshnev S.V. Gibridnaja zhidkostnaja model' magistral'nogo Internet-kanala/ S.V. Porshnev, M.K. Grebenkin// Saarbrücken: LAP LAMBERT Academic Publishing GmbH & Co. KG, 2012. 163 p.
7. Grebenkin M.K., Porshnev S.V. Programmnaja realizacija gibridnoj zhidkostnoj modeli informacionnyh potokov v vysokoskorostnyh magistral'nyh internet-kanalah// Svidetel'stvo o gosudarstvennoj registracii programm dlja JeVM № 2012616118 (Zajavka no. 2012613786. Data postuplenija 11 maja 2012 g. Data registracii v Reestre programm dlja JeVM 4 ijulja 2012 g.)
8. Setevoy simuljator OPNET, Available at: <http://opnet.org> (accessed at 18.07.2013).
9. Hinchin, A.Ja. Matematicheskie metody teorii massovogo obsluzhivaniya/ A.Ja. Hinchin // Tr. Mat. Instituta im. V.A. Steklova AN SSSR. 1955.
11. D-ITG, Distributed Internet Traffic Generator, Available at: <http://www.grid.unina.it/software/ITG/download.php> (accessed at 18.07.2013).
12. Marsan M.A., Garetto M, Giaccone P., Leonardi E., Schiattarella E., Tarello A. Using Partial Differential Equations to Model TCP Mice and Elephants in Large IP Networks. IN: FOCOM 2004. Twenty-third Annual Joint Conference of the IEEE Computer and Communications Societies. 2004. Vol. 4. pp. 2821–2832.
13. Misra, V., W.-B. Gong, Towsley D. Fluid-Based Analysis of a Network of AQM Routers Supporting TCP Flows with an Application to RED. ACM SIGCOMM Computer Communication Review. 2000. Vol. 30. Issue 4. pp. 151–160.
14. Tmix: A Tool for Generating Realistic TCP Application Workloads in ns-2, Available at: <http://www.sigcomm.org/ccr/drupal/files/p67-v36n3l-weigle.pdf> (accessed at 18.07.2013).
15. Xu Zh., Bagrodia R. GPU-Accelerated Evaluation Platform for High Fidelity Network Modeling, Proceeding PADS '07 Proceedings of the 21st International Workshop on Principles of Advanced and Distributed Simulation. 2007. pp. 131–140.

Рецензенты:

Доросинский Л.Г., д.т.н., профессор, зав. кафедрой информационных технологий, ГАОУ ВПО «Уральский федеральный университет им. первого Президента России Б.Н. Ельцина», г. Екатеринбург;

Суханов В.И., д.т.н., доцент, зав. кафедрой микропроцессорной техники, ГАОУ ВПО «Уральский федеральный университет им. первого Президента России Б.Н. Ельцина», г. Екатеринбург.

Работа поступила в редакцию 08.10.2013.