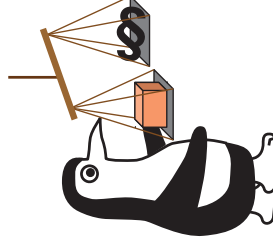




Die Deutsche Post ist besser als ihr Ruf oder für die echten Fans: es gibt nur einen Proof. Bungartz, nur einen Proof. Buuungartz. ...

Inhalt

Editorial	2
Performance Engineering in Aktion	4
Iterationsschleife	7
SeisSol: Erdbeben-Simulationen auf Petascale	8
Lattice-Boltzmann-Löser im Verbundprojekt	12
Absolventenfeier des ENB 2013	16
Im Endspurt: KAUST Projekt Virtual Arabia	18
SPPEXA News	20
5 weitere Jahre BGCE!	21
ASCETE Sudelfeld Summit	22
Workshop on ICCE 2014	22



Das Quartl erhalten Sie online unter <http://www5.in.tum.de/quartl/>



Das Quartl ist das offizielle Mitteilungsblatt des Kompetenznetzwerks für Technisch-Wissenschaftliches Hoch- und Höchstleistungsrechnen in Bayern (KONWIHR) und der Bavarian Graduate School of Computational Engineering (BGCE)

Editorial

Die Weihnachtszeit ist vorbei – selbst die Hartnäckigsten unter den Weihnachtsbeleuchtungsfreaks (kennen Sie die Familie Griswold?) haben an Mariä Lichtmess den Schalter betätigt und die Lichtanlage abgebaut. Dennoch wollen wir nochmals in den Advent zurückblicken, um die geradezu dramatischen Auswirkungen näher zu betrachten, die eine einfache Dose mit Lebkuchen mit sich bringen kann.

Was war geschehen? Nun, zu Beginn der Adventszeit erreichte mich ein mittelgroßes Päckchen – und dieses gelangte zu mir, obwohl es die höchst merkwürdige Anschrift

Ludwig-Maximilians-Universität München¹

Prof. Hans-Joachim Bungartz

Boltzmannstraße 3, 85748 Garching

trug. Fast alles richtig, bis auf die Uni halt. Auch wenn sich mir die bisweilen an Absurdität kaum zu toppende Beziehungskiste zwischen den beiden großen Münchner Unis noch nie wirklich erschlossen hat, so war hier doch Vorsicht geboten. Was würde da wohl drin sein – am Ende eine Briefbombe? Und wer verschickt überhaupt Päckchen mit solchen Adressen? Also rasch den Absender inspiziert – aha, ein in IT-Kreisen nicht ganz unbekanntes Unternehmen aus der Magenta-Ecke. Hmm, eigentlich kennt man sich ja schon; vielleicht sollte ich demnächst mal Magenta-Post nach Bonn schicken und an Vodafone adressieren, oder so ähnlich. Aber halt, das sind dann doch zu niedere Rachegeleüste. Das eigentlich Beunruhigende an der Sache (insbesondere für alle Kunden der Lebkuchen-versendenden Firma) ist, dass derartige Päckchen heutzutage ja nicht mehr manuell von Menschen adressiert werden; vielmehr hat die IT hier durchgängig das Sagen. Und dass die gerade bei einem IT-Konzern offensichtlich nicht so funktioniert, wie sie sollte, grenzt schon ans Peinliche. Aber die Jungs haben schon mehrfach büßen müssen: Bei der Mitgliederversammlung des DFN Anfang Dezember war das natürlich auch eine willkommene Steilvorlage für einen Gag.

Doch wenden wir uns nun dem Inhalt des Päckchens zu – einem der unumstrittenen fränkischen Exportschlager, eine Dose mit Nürnberger Lebkuchen.

¹ siehe Seite 24

simulation with the enabling disciplines computer science and mathematics. Thus, in particular young researchers and PhD students are encouraged to organize a minisymposium with 4-5 speakers focussed on a topic related to their thesis or project work or to an important enabling methodology or submit regular contributions.

Motivated by current trends in application fields, new requirements arising from rapidly increasing parallelism in compute architectures, and novel mathematical approaches, the workshop focusses in particular on the subtopics

- scalable parallel algorithms
- simulation-based optimization
- uncertainty quantification.

On Monday, October 6, and Tuesday, October 7, we offer several tutorials presenting the state-of-the-art in selected workshop topics and enabling the participants to actively use the presented methods and tools in their research projects.

For more details, see

<http://ipvs.informatik.uni-stuttgart.de/SGS/ICCE2014/index.php>

M. Mehl

Quartl* - Impressum

Herausgeber:

Prof. Dr. A. Bode, Prof. Dr. H.-J. Bungartz, Prof. Dr. U. Rüde

Redaktion:

J. Daniel, C. Halfar, C. Kowitz, Dr. S. Zimmer

Technische Universität München, Fakultät für Informatik

Boltzmannstr. 3, 85748 Garching b. München

Tel./Fax: ++49-89-289 18630/18607

e-mail: halfar@in.tum.de, **www:** <http://www5.in.tum.de/quartl>

Redaktionsschluss für die nächste Ausgabe: **31.05.2014**

* Quartl: früheres bayerisches Flüssigkeitsmaß,

→ das Quart: 1/4 Kanne = 0.27 l

(Brockhaus Enzyklopädie 1972)

Auch an dieser Stelle nochmals herzlichen Dank an alle Beteiligten, die in irgendeiner Form zur erfolgreichen Weiterführung beigetragen haben – insbesondere geht der Dank natürlich an die Trägeruniversitäten FAU Erlangen und TU München sowie die drei involvierten Fakultäten, die sich bereitklärt haben, die entsprechende Finanzierung zu übernehmen.

T. Neckel

Workshop zur numerischen Simulation von Erdbeben und Tsunamis: ASCETE Sudelfeld Summit

Seit 2 Jahren wird innerhalb des ASCETE Projektes an der numerischen Simulation von Erdbeben und Tsunamis geforscht.

Gegen Ende des von der Volkswagen Stiftung geförderten Projektes wird nun ein Workshop organisiert auf dem mit führenden Wissenschaftlern in den Gebieten Tsunamis, Erdbeben, numerische Methoden und HPC der aktuellen Forschungsstand diskutiert werden soll. Wir versprechen uns viele neue Einblicke durch die Interdisziplinität des Treffens in einer gemütlichen Atmosphäre in den Bayrischen Alpen bei Bayrischzell. Natürlich werden wir ausführlich im Sommer-Quartl berichten, wie der am 20. – 23.5. stattfindende Workshop gelaufen ist. Informationen zum Workshop und zur

Anmeldung finden sich auf www.ascete.de

M. Bader, J. Behrens, C. Pelties

The 3rd International Workshop on Computational Engineering ICCE 2014

The 3rd International Workshop on Computational Engineering ICCE 2014 is going to be held in Stuttgart from October 6 – October 10, 2014. The first two International Workshops on Computational Engineering have been held in Munich (2009) and Darmstadt (2011).

The 3rd International Workshop on Computational Engineering in Stuttgart aims at bringing together competences in application fields of numerical

chen. Sofort schossen mir diverse Gedanken durch den Kopf: War das endlich mal ein Bestechungsversuch, würde mich der genussvolle Biss in einen Lebkuchen zum korrupten Beamten machen? Und bestimmt würde es in diesen Compliance-durchdrungenen Zeiten irgendwo ein dickes Regelwerk geben, wie denn nun korrekt mit einem derart unmoralischen Angebot umzugehen sei. So was in der Art „Bis 2 Euro fuffzig wird es toleriert, danach heißt es dann 'Ab in den Bau!'“; oder die Pension geht flöten, oder was auch immer. Zurückschicken? Dann müsste aber bitte fairerweise die LMU das Porto übernehmen. Oder etwa bei der Steuer als Einkünfte anmelden? Denn wer weiß, wie lange es die Option der Strafbefreiung per Selbstanzeige später noch geben wird. Aber was für Einkünfte sind das denn überhaupt, schließlich habe ich ja kein Gewerbe angemeldet, das auf Einnahmen in Form von Virtualien abzielt.

Vergammeln lassen? Hilft nicht wirklich, und wäre ja schon schade. Also am Ende doch verzehren, einfach so? Aber dann vorher wenigstens die Schuld auf verschiedene Schultern verteilen, sozusagen eine Kollektivschuld generieren (obwohl es die ja laut einem unserer früheren Bundespräsidenten gar nicht gibt – schon grundsätzlich nicht, also insbesondere auch bei Lebkuchenkonsument nicht)?

Nun, in den Regularien habe ich nicht nachgeschaut, auch habe ich unsere Rechtsabteilung nicht konsultiert – schließlich haben Lebkuchen als Lebensmittel ja ein Verfallsdatum. Somit haben wir das am Lehrstuhl gemeinsam verfütert. Ein kleines maues Gefühl hatte ich dabei im Magen schon, aber geschmeckt hat mein Anteil dennoch. Trotz der skandalösen Adresse. Und sollte demnächst eine Rücktrittsforderung für eines meiner Ämter eintreffen – wer weiß, vielleicht wäre ich über diese Wendung in der „Lebkuchen-Affäre“ ja gar nicht so unglücklich?

Mit der Compliance-Thematik ist es wie (leider) so oft: Eine wichtige Sache wird derart übertrieben reguliert, bis in jedes popelige Detail hinein, dass am Ende oft Skurriles steht; das darüber hinaus dann auch noch oft auf höchst skurrile Weise interpretiert wird. Z.B., wenn bei Firmen Lebkuchendosen tatsächlich zurückgeschickt werden, aus Sorge vor rechtlichen Implikationen; oder wenn in einer Behörde (Name dem Verfasser bekannt) gemäß einer hausinternen Festlegung besonders strenge Regeln für alle Besoldungsgruppen bis E14/A14 („bis“ von unten zu interpretieren) gelten –

oberhalb sind die Damen und Herren natürlich moralisch so gefestigt, dass nichts zu fürchten ist und man es folglich nicht so genau nehmen muss.

Doch genug gelästert – die gesamte Quartl-Redaktion hofft, dass Sie gut ins Neue Jahr gestartet sind und auch die närrische Zeit ohne Blessuren überstehen. In diesem Sinne viel Spaß mit der neuen Ausgabe Ihres Quartls und eine gute Zeit!

H.-J. Bungartz

Performance Engineering in Aktion

Am 3. und 4. Dezember 2013 fand am LRZ zum zweiten Mal der zweitägige PRACE²-Kompaktkurs „Node-Level Performance Engineering“ statt. Georg Hager und Gerhard Wellen vom RRZE vermittelten den über fünfzig Teilnehmern Grundlagen der Prozessor- und Rechnerarchitektur und der modellbasierten Performance-Optimierung auf dem Rechenknoten. Das Gelernte konnte an einem FE-Code des Lehrstuhls für Computation in Engineering der TUM gleich in die Praxis umgesetzt werden.

Die Vortragenden spannten einen Bogen von der Architektur moderner Multicore-Prozessoren und -Systeme über Mikrobenchmarking und das Roofline-Modell bis zur modellgesteuerten Code-Optimierung. Dabei wurde besonderer Wert auf die Frage gelegt, wie man die „Lichtgeschwindigkeit“ einer Programmschleife bestimmt, d.h. die maximale Performance, mit der sie auf einem Prozessor laufen kann. Davon ausgehend ergeben sich dann meist Vorhersagen über die Konsequenzen bestimmter Code-Transformationen. Anstatt also auf gut Glück verschiedene Lehrbuch-Optimierungen auszuprobieren, arbeitet man sich in strukturierter Weise an einen „optimalen“ Code heran. Nach einer Vertiefung zu den wichtigen Themen SIMD, ccNUMA und SMT schloss der Kurs mit einer Darstellung des ECM-Modells³ ab, das als Verfeinerung des bekannten Roofline-Modells auch die Vorhersage

²PRACE (Partnership for Advanced Computing in Europe)

<http://www.prace-project.eu/>

³Details finden sich z.B. in <http://dx.doi.org/10.1002/cpe.3180>

Das Projekt EXA-DUNE bietet vom 24. bis 28. März einen Intensivkurs zur Software DUNE in Heidelberg an.

Zum Monatswechsel sind ein Workshop zum Projekt ExaStencils (31. März - 1. April) sowie ein Gender Training in Dresden angesagt.

Anschließend an den Workshop und das Gender Training wird von der TU Dresden die diesjährige Jahresvollversammlung von SPPEXA am 2. April ausgerichtet. Bei dem Zusammentreffen werden sowohl die Fortschritte aller beteiligten Projekte als auch mögliche Kooperationen und die zukünftige Planung des Schwerpunktprogramms thematisiert.

Auch 2014 wird ähnlich zum Vorjahr wird ein SPPEXA-Doktorandenkurs („Doctoral Retreat & Coding Week“) stattfinden. Der Kurs mit vorwiegend „Hands-on“-Charakter wird im Rahmen der Ferienakademie vom 21. September bis 3. Oktober veranstaltet und von Philipp Neumann, Josef Weidenborfer (beide TU München) und Karl Frlinger (LMU München) organisiert. Neben zwei eingeladenen Tutorials wird der diesjährige Schwerpunkt auf effizienten Algorithmen und hardwarenaher Programmierung liegen.

P. Neumann

5 weitere Jahre BGCE!

Die Erfolgsgeschichte geht weiter: Nach der Begutachtung der Bavarian Graduate School of Computational Engineering (BGCE) am 19. November 2013 hat die Internationale Kommission des Elitenetzwerks Bayern die Weiterführung der BGCE Anfang Januar nun einstimmig bewilligt.

Damit sind weitere fünf Jahre dieses Elitestudiengangs gesichert, bis die nächste Evaluation durch das Elitenetzwerk und entsprechende Fachgutachter dann voraussichtlich im Jahr 2019 stattfinden wird. Das ist vor allem für die potentiellen neuen Studierenden aus den drei Basisstudiengängen Computational Engineering (CE), Computational Mechanics (come) und Computational Science and Engineering (CSE) eine prima Nachricht, die nun auch zukünftig einen Master of Science with honours erwerben und vor allem an diversen interessanten Kursen und Angeboten des Programms teilnehmen können.

Projekthema extern vorgegeben war, so haben doch alle darin die nötige Freiheit gefunden, um sich wissenschaftliche zu verwirklichen.

Und so blicken wir zurück mit einem lachenden und einem weinenden Auge: Lachend, weil das Projekt erfolgreich und mit Highlights gespickt war (man denke nur an die gemeinsamen mehrmonatigen Auslandsaufenthalte, gemeinsame Konferenzbesuche und Publikationen, die fertigen oder fast fertigen Dissertationen, ...). Weinend, weil es nun aus ist.

T. Weinzierl

P.S. Wie im internationalen Film üblich, kommt jetzt eine Liste an Haupt- und Nebenrollen. Überhaupt nicht repräsentativ oder vollständig.

Hans-Joachim Bungartz	Oberindianerhüpfeling
N.A.	Antagonist (bis heute nicht bekannt)
Scherge 1	die Juristen
Scherge 2	der/die Änderer aller Report-Sheets und Formulare
Scherge 3	der das Visum ausstellt und die Unterkünfte organisiert
Marcus Tönnis	Mr. Q (oder der mit den Equipmentideen)
Wolfgang Eckhardt	the Replacement-Killer
Stefan Zimmer	Scherge 1-Töter
Joelle Coutinho	Scherge 2-Töter
Jens Schneider	Scherge 3-Töter



Im Frühjahr 2014 haben im Rahmen des DFG Schwerpunktprogramms SPPEXA⁵ erneut verschiedene Aktivitäten stattgefunden oder sind geplant.

Auf der SIAM PP (SIAM Conference on Parallel Processing for Scientific Computing) in Portland, Oregon, war SPPEXA durch zwei Minisymposia „Towards Multilevel Solvers for Exascale“ und „Implementation Aspects of Parallel-in-time Methods on HPC Systems“ vertreten.

Ein weiteres Young Researcher’s Minisymposium zu „Time-parallel methods“ hat im Rahmen des GAMM Annual Meetings am 11. März in Erlangen stattgefunden.

⁵Software for Exascale Computing

der Einzelkernperformance und Skalierung einer parallelen Schleife auf einem Multicore-Prozessor erlaubt.

Am Lehrstuhl für Computation in Engineering wurde sogleich versucht, diese Konzepte auf den am Lehrstuhl entwickelten Code für Finite Elemente hoher Ordnung anzuwenden. Wie bei FE-Programmen typisch, befasst sich dabei ein wesentlicher Teil des Codes mit der numerischen Integration der Steifigkeitsmatrix K . Vereinfacht betrachtet wird diese Aufgabe in zwei ineinander geschachtelte Abschnitte aufgeteilt: Eine äußere Schleife iteriert über alle Elemente im FE-Netz und eine innere Schleife läuft über alle Integrationspunkte des jeweiligen Elements. Dabei muss auf jedem Integrationspunkt die Elementsteifigkeitsmatrix wie folgt integriert werden:

$$K_e = K_e + B^T C B.$$

Hier ist B der diskrete Dehnungsoperator und C die Materialmatrix. Die integrierte Elementsteifigkeitsmatrix wird dann in einem weiteren Schritt in die globale Steifigkeitsmatrix assembliert. In Pseudocode sieht dieser Programmteil wie folgt aus:

```
1 K = 0;
2 for all elements
3 {
4   Ke = 0;
5   for all integration points of current element
6     // skaliert mit O(p^3)
7   {
8     B = computeB( currentIntegrationPoint )
9     tmpProduct = trans(B) * C;
10    Ke += tmpProduct * B;
11    // skaliert mit O(p^6)
12  }
13  assembleElementStiffnessMatrix( Ke, K );
14 }
```

Ein wichtiger Unterschied zwischen Finiten Elementen niedriger und hoher Ordnung ist nun die Lastverteilung auf die verschiedenen Loop-Level. Werden nur lineare Ansatzfunktionen verwendet ($p = 1$), ist das Verhältnis der benötigten Integrationspunkte zur Anzahl der erforderlichen Elemente

klein. Darüber hinaus ist auch die Matrix B klein, sodass die innere Schleife wenig Last verursacht. Damit liegt die Last auf der äußeren Schleife. Da diese Form der Finiten Elemente stark verbreitet ist, sind die meisten Parallelisierungs-Frameworks auf diesen äußeren Loop hin optimiert (z.B. PETSc und Trilinos).

Im Falle von Finiten Elementen hoher Ordnung ($p \geq 3$), ist die Lastverteilung gespiegelt. Es werden wenige Elemente eingesetzt, auf denen aber Polynome hoher Ordnung als Ansatzfunktionen aufgespannt werden. Dies erhöht die Anzahl der Integrationspunkte und die Größe der B Matrix gleichermaßen, was die Hauptlast auf die innere Schleife verlagert.

Für hohe Ordnungen werden daher bis zu 60% der gesamten Programmaufzeit für die Multiplikation in Zeile 9 benötigt. Jede Verbesserung dieses Flaschenhalses hat daher einen signifikanten Einfluss auf die Gesamtpformance des Codes. Zu diesem Zweck wurde dieser Teil bereits zu einem früheren Zeitpunkt als Fortranroutine in den C++ Code eingebaut, was zu einer beachtlichen Steigerung der Leistung führte.

Motiviert durch den oben angesprochenen Kurs wurde dieser Programmaschnitt dann in Zusammenarbeit mit der HPC-Gruppe des RRZE verfeinert. Hierbei stellte sich heraus, dass es sich um ein typisches Beispiel von „Compiler-Psychologie“ handelte – wie bringt man den Compiler dazu, den Code zu erzeugen, der die optimale Performance garantiert? Auch mit Fortran ist das oft nicht so einfach und erforderte in unserem Fall einige manuelle Eingriffe. Das beliebte Paradigma „Der Compiler wird’s schon richten“ war hier ganz klar fehl am Platz! Dank eines strukturierten Performance-Engineerings (wie im Kurs gelernt) konnte schließlich über 80% der vom ECM-Modell vorhergesagten Einzelkern-Performance erreicht werden, mit einer fast perfekten Skalierung über die Kerne in einem Sockel. Insgesamt ergab sich im Vergleich zum Originalprogramm eine Beschleunigung um einen Faktor zwei.

Die zwei Tage am LRZ waren in jedem Fall gut investierte Zeit. Wir hoffen auch in Zukunft mehr von diesem Wissen in die Code-Entwicklung am Lehrstuhl einbringen zu können und freuen uns über die gelungene Zusammenarbeit.

T. Bog, J. Frisch, G. Hager, S. Kollmannsberger, E. Rank, N. Zander

virtuellen arabischen Lichter aus. Nicht, ohne eine eintägige Versammlung des gesamten Teams, um noch einmal auf das Vergangene und auch Vollbrachte zurück zu schauen.



Abbildung 7: Computational Scientists@Tent/Work

Dabei wurde anhand des Münchner Olympiastadions nochmals ausgiebig das Prinzip der (Beduinen-) Zeltkonstruktion bestaunt, Pizza in großen Mengen vernichtet, und sich gegenseitig Feedback zum individuellen Projektbeitrag gegeben. Den Tag abgerundet hat dann ein gemütliches Abendessen. Beim konstruktiven Feedback wurden einige Dinge deutlich: Auch wenn zumeist die Kollaboration nur P2P erfolgt ist, war die transitive Hülle des Kollaborationsgraphen dann doch eine Clique – das Projekt hat mustergültig zusammengearbeitet. Auch wenn die Arbeit an brandneuem Equipment als Forschungsgegenstand an sich allen Doktoranden Spaß gemacht hat, es hat halt auch eine Menge Zeit gekostet – aber das Projekt hat auch etwas herausragendes physisch Vorzeigbares an der TUM hinterlassen. Auch wenn das

was ich mir vorgenommen hatte.

Ich hoffte schon, schlimmer kann's nimmer werden. Denkste, erwiesen sich die beiden Vorträge der Absolventen als flotte, motivierende und überaus interessante Präsentationen. Wo waren all die schönen, überfordernden Formeln geblieben? Wo die 100 Folien ohne ein einziges Bild? Des hamma nämlich scho ollawei so g'macht. Anstatt grauer Theorie gab's stattdessen vorgeführte Praxis, die sogar noch funktionierte. Ein Graus! Aber irgendwie musste sich dieser Abend doch in die Länge ziehen lassen. Natürlich: Bei der Einzelzelehrung der Absolventen, da witterte ich die letzte große Chance auf einen langweiligen Abend: Name vorlesen, aus dem Publikum nach vorne stolpern, Zeugnis überreicht bekommen, das braucht alles Zeit. Jedoch, die geneigte Leserin/der geneigte Leser ahnt es schon: Auch hier machten Effizienz und Planung meine Erwartungen zu Nichte. Im 7-Sekunden-Takt eine respektvolle aber prompte Ehrung nach der anderen.

Das i-Tüpfelchen der Enttäuschung waren dann die kreativen Filmenspieler, die von der Uni Ansbach produziert wurden. Fakten, Fakten, Fakten für jedermann verständlich aufbereitet in bewegtem Bild und Ton. Mein Langweilenickerchen konnte ich mir endgültig abschreiben.

Dass es zum Ausklang des Abends Sekt, variationsreiche Schnittchen und piffige Knödel in Pilzsoße anstatt des in Bayern üblichen Bier und Brez'n gab, schockierte mich dann auch nicht mehr. Wos glaum's denn, wos des kost'?

Somit lautet mein Fazit: Noch nie wurden meine Erwartungen und Hoffnungen auf eine langweilige, verstaubte, träge, sich in die Länge ziehende, schlecht organisierte und undurchdachte Absolventenfeier so massiv enttäuscht wie im Jahr 2013 vom ENB.

C. Riesinger

Es ist vollbracht: KAUST Projekt Virtual Arabia

Fünf Jahre lang ist das KAUST Projekt Virtual Arabia an der Technischen Universität München gelaufen und fünf Jahre lang war es treuer Begleiter des Quartl.

Wie jedes Projekt, musste nun aber auch dieses Projekt einmal zu Ende gehen und so gingen am 31. Dezember mit den Sylvesterböllern auch die

Iterationsschleife N=11

17. Februar 2014

Der italienische Schriftsteller Carlo Lorenzini wurde am 24. November 1826 in Florenz geboren, das damals noch von Habsburgern regiert wurde (die Medici-Linie war 1737 ausgestorben). Außer als Schriftsteller verdiente Lorenzini seinen Unterhalt als Journalist und als Angestellter der Präfektur. Er soll sich an den italienischen Unabhängigkeitskriegen beteiligt haben, in deren Verlauf Florenz kurzzeitig italienische Hauptstadt wurde (1864 – 1870), sodass er am 26. Oktober 1890 in seiner Heimatstadt in einem italienischen Nationalstaat sterben konnte^a.

Carlo Lorenzini ist weder als Unabhängigkeitskämpfer, noch als Journalist noch als Präfekturangestellter in Erinnerung geblieben. Dass es noch immer Spuren von ihm gibt, verdankt er vielmehr einem Buch, das er in den Jahren 1881 bis 1883 geschrieben hat. Unter dem Pseudonym Carlo Collodi begann er 1881 in der Zeitschrift „Giornale per i bambini“ eine Fortsetzungsgeschichte über einen Hapelmann. Der Erfolg der Fortsetzungsgeschichte war so groß, dass daraus eine Buchidee entwickelt wurde, die im Februar 1883 ihre erste Auflage erlebte. Der Name des Hapelmanns war Pinocchio, und das Buch hat seither weltweit Verbreitung gefunden und wurde vielfach verfilmt (zuerst vermutlich von einem E. Pasquali, über den genauere Information zu finden dem Autor der Iterationsschleife nicht gelang).

Das Buch selbst schockiert beim ersten Lesen durch seine Skurrilität und Herzlosigkeit. Ein alleinerziehender älterer Mann kümmert sich zwar kindlich rührend um seinen im wahrsten Sinne des Wortes selbst-geschnitzten Sohn, vernachlässigt aber sträflich seine Aufsichtspflichten und setzt das Kind damit einer Reihe von Gefahren aus, denen Pinocchio – wie zu erwarten – nicht gewachsen ist. Der hilflose Versuch, dem Kind Halt im Leben durch die moralischen Ratschläge einer Grille zu geben, scheitert logischerweise ebenso kläglich wie das naive Experiment, das Kind durch Schulbesuch zu moralisch-ethischem Verhalten zu bringen.

^aNoch heute findet man in einer kleinen Seitengasse in Florenz das Büro der „Associazione Nazionale Veterani e Reduci Garibaldini“. Was zunächst als rein historische erscheint, entpuppt sich als real existierender Verein mit Webseite <http://nuke.garibaldini.com/> und einer regelmäßig erscheinenden Zeitschrift. Neben aktuellen Berichten über italienische Soldaten in Afghanistan finden sich historische Artikel und Erinnerungen an verstorbene Kollegen. Die ANVRG scheint eine Art Kameradschaftsbund italienischer Prägung zu sein, der eine Linie von 1864 bis heute zu ziehen versucht.

So stürzt Pinocchio von einer Malaise in die andere, wird von Katern, Füchsen, Eseln und Zirkusdirektoren über den Tisch gezogen wie ein Jungendlicher bei Facebook und landet zuletzt im Bauch eines Walfisches. Erlösung bringt zuletzt die Erkenntnis, dass mit dem verkorksten Leben eines arbeitslosen Tischlers kein Staat zu machen ist, und die Aneignung eines Häuschens im Grünen mit liebender Fee als Mutter der einzige Weg zum erfolgreichen Eintritt in die Erwachsenenwelt sein kann. Versteckt in all dem skurrilen Gewirr von Geschichten findet sich jedoch eine kleine Weisheit: Pinocchio wird nicht menschlich, weil er brav geworden ist, sondern weil er eine sinnvolle Balance zwischen seinen eigenen Ansprüchen an das Leben und der Rücksicht auf seine Mitmenschen findet.

Indem Pinocchio Gier und Neid des Kindes hinter sich lässt, wird er erwachsen und erst damit Mensch. Was Collodi in einem Kinderbuch verstecken wollte – oder musste (?) – ist noch heute schwer zu finden, aber die Mühe, Pinocchio noch einmal zu lesen, lohnt.

M. Resch

SeisSol: Erdbeben-Simulationen auf Petascale

Erdbebenprozesse besser verstehen und genauer quantifizieren zu können ist eine der „Grand Challenges“ des Wissenschaftlichen Rechnens und der Seismologie.

Um die entsprechenden Multiphysik-Probleme – vom nichtlinearen Bruchprozess entlang einer Verwerfung bis zur Ausbreitung der angeregten seismischen Wellen – für realistische 3D-Erdmodelle mit heterogenen Materialeigenschaften der Erdkruste in der notwendigen Auflösung simulieren zu können, sind nicht nur entsprechend leistungsfähige Supercomputer notwendig, sondern auch entsprechend effiziente Simulationssoftware.

SeisSol⁴ ist einer der wenigen Simulationscodes, die Erdbebenbruchdynamik gekoppelt mit seismischer Wellenausbreitung akkurat berechnen können. SeisSol löst die elastischen Wellengleichungen mit dem Arbitrary high-order DERivatives Discontinuous Galerkin (ADER-DG) Verfahren, das eine hohe Approximationordnung in Raum und Zeit ermöglicht.

⁴ <http://seissol.geophysik.uni-muenchen.de/>

mit Glanz und Glamour. Das war ja mal so gar nicht das, was ich erhofft hatte. Des hod's ja no nia ned ge'm! Aber der unheilvolle Abend hatte erst begonnen.

Denn wenn es schon unbedingt ein Theater sein musste, dann doch immerhin mit der passenden musikalischen Untermalung. Klassische Musik war sozusagen Pflicht. Stattdessen wurden einem Pop-Songs um die Ohren gehauen und Musical-Klassiker aufgeführt. Fast könnte man meinen, die Organisatoren wären sich des jungen Publikums bewusst gewesen. Und die Versuchung war groß, die Songs mitzusingen, die da gerade performt wurden. Wo kemma denn do hi?



Abbildung 6: Gruppe der geehrten BGCE-Absolventen/Absolventinnen mit Staatssekretär Bernd Sibler (v.l.)

Kaum von diesem Schock erholt, ging es weiter mit enttäuschter Vorfreude auf einen eingetasteten Abend. Immerhin hatte man mit Herrn Bernd Sibler den Bayerischen Staatssekretär für Wissenschaft und Kunst gewinnen können. Anstatt der erhofften Ähm's und Ähnm's und hölzerner Reden leitete Herr Sibler zu meiner Überraschung flott mit Witz und Spontanität durch den Abend. Ständig war man verleitet, mitzulachen. So gar nicht das,

Absolventenfeier des ENB 2013

(K)ein Hort von Langeweile

Am 2. Dezember lud das Elitenetzwerk Bayern (ENB) zu seiner alljährlichen Absolventenfeier, bei der diejenigen Personen geehrt wurden, die im vergangenen Jahr ihr Studium in einem der 21 Elitestudiengänge oder ihre Promotion in einem der 14 internationalen Doktorandenkollegs abgeschlossen haben. Ebenfalls wurden die Mitglieder des Max Weber-Programms und der Nachwuchsforschergruppen und Inhaber von Forschungsstipendien ausgezeichnet. Als Programmkoordinator eines Elitestudiengangs des ENB, der BGCE, ist das natürlich eines der Highlights des Jahres, wird doch bei einem solchen Ereignis alles aufgeföhren, was dem verstaubten akademischen Klischee entspricht. Doch dieses Jahr sollte ich herb enttäuscht werden.



Abbildung 5: Blick auf die Ränge des Prinzregententheaters mit Organisatoren, Absolventen, Koordinatoren und Gästen.

Ich hatte mich schon darauf gefreut, den Abend in einem alten, muffigen Hörsaal zu verbringen. Kaputte Klappstühle und eine Atmosphäre zwischen Prüfungsangst und Professorendrill sollten mich wehmütig an meine eigene Studentenzeit zurück erinnern. Und dann das: Stattdessen fand die Feier im Prinzregententheater in München statt, einer top-schicken Kultureinrichtung

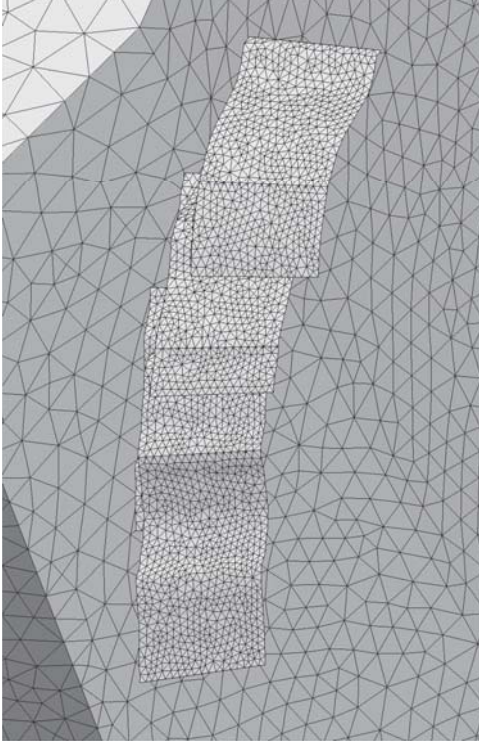


Abbildung 1: Beispiel für ein kompliziertes System geologischer Verwerfungen (Bruchsegmente des Magnitude 7.3 Landers-Erdbebens, 1992).

Die nichtlinearen Bruchprozesse auf der zweidimensionalen Erdoberfläche dienen als Quelle für die Wellenausbreitung, wobei auch Rückkopplungen der Wellenausbreitung auf den Bruchprozess berücksichtigt werden. Eine solche Bruchebene, die unter Umständen sehr komplex sein kann, ist in Abbildung 1 illustriert. In einem gemeinschaftlichen Projekt der Arbeitsgruppe Geophysik der LMU München, (Dr. Alice-Agnes Gabriel, Dr. Christian Pelties, Stefan Wenk) und der HPC-Arbeitsgruppe von Prof. Michael Bader (Alexander Breuer, Alexander Heinecke, Sebastian Rettenberger) wurde SeisSol in den vergangenen anderthalb Jahren für Petascale-Simulationen fit gemacht: als krönender Abschluss konnte Ende Januar eine Wellenausbreitungssimulation im Vulkan Merapi mit 100 Millionen Gitterzellen und fast 50 Milliarden Freiheitsgraden auf dem SuperMUC durchgeführt werden, wobei die „magische“ Marke von 1 PetaFlop/s (10^{15} Rechenoperationen pro Sekunde) erreicht wurde.

Gefördert wurden die Arbeiten unter anderem durch die DFG, ein entscheidendes KONWIHR-Projekt und das von der Volkswagen Stiftung geförder-

te Projekt ASCETE (Advanced Simulation of Coupled Earthquake-Tsunami Events).

Optimierung von SeisSol für PetaScale-Simulationen

An drei wesentlichen „Baustellen“ wurde SeisSol zu diesem Zweck entscheidend verbessert:

Optimierte Matrix-Kernel: In SeisSols ADER-DG-Methode werden die innersten, Zell-lokalen numerischen Operationen über Matrixmultiplikationen implementiert. Für diese wurden über einen Codegenerator-Ansatz hochoptimierte, Hardware-spezifische Matrix-Kernel eingesetzt. Abhängig von der Besetztheitsstruktur der jeweiligen Matrizen werden Kernels für dicht- oder dünnbesetzte Matrizen gewählt – je nach kürzerer erzielter Laufzeit. Durch die bessere Ausnutzung von Vektoreinheiten (und anderer Hardware-Eigenschaften) aktueller CPUs wird dadurch (je nach Approximationsordnung) eine Beschleunigung um einen Faktor 5 und mehr erreicht.

Hybride MPI+OpenMP-Parallelisierung Die Parallelisierung von SeisSol wurde konsequent auf eine hybride MPI+OpenMP-Parallelisierung umgestellt. Die entsprechende Reduzierung der MPI-Prozesse wirkte sich, insbesondere bei einer sehr großen Zahl von verwendeten Rechenkernen, positiv auf die parallele Effizienz aus. Insbesondere wurden auch Management- und Kommunikationsroutinen parallelisiert, die ebenfalls bei großen Core-Zahlen die Performance beeinträchtigen.

Paralleles Einlesen der Gitter SeisSols Stärken basieren auf der Verwendung beliebig adaptiver Tetraedergitter, die notwendig sind um auch komplizierteste Geometrien akkurat modellieren zu können. Diese werden über einen externen Gittergenerator (z.B. SimModeler von Simmatrix) erzeugt. Von Tetraedergittern ist ein weiterer Vorteil die schnelle und teilweise automatisierbare Generierung größter Gitter mit minimaler Nutzerinteraktion. So ist eine minimale “time-to-solution” sichergestellt. Als wesentlicher Schritt in Richtung Petascale-Simulationen wurde das Einlesen dieser Gitter auf Basis von MPI I/O neu implementiert. Dadurch können nun auch Gitter mit mehr als einer Milliarde Gitterzellen auf über tausenden von Rechenknoten effizient eingelesen werden.

zelter Teilschritte des Lattice-Boltzmann-Lösers erreicht. Diese beiden Optimierungen sorgten für eine fünffache Beschleunigung. Die Performance konnte zusätzlich gesteigert werden durch die Nutzung von nativen C-Arrays statt C++-STL-Containern zur Repräsentation des Strömungsfelds. Die C-Arrays bieten zwar nicht die komfortablen Möglichkeiten des Zugriffs und der Überprüfung auf Einhaltung der Array-Grenzen eines STL-Containers, sind aber für den Compiler deutlich einfacher zu optimieren.

Um die gute Parallelisierbarkeit der Lattice-Boltzmann-Methode für Mehrprozessorsysteme zu nutzen, wurde eine OpenMP-Parallelisierung durchgeführt. Diese sorgt dafür, dass die Performance nahezu mit der Anzahl der Prozessorkerne skaliert, bis die Teilgebiete zu klein werden und schließlich der Overhead der Synchronisation dominiert. Die gute Skalierung ist möglich, weil relevante Problemgrößen im Cache der modernen Intel-Prozessoren gehalten werden können und somit die Hauptspeicher-Bandbreite keinen Flaschenhals darstellt.

Betrachtet man die Performance in GFLOP/s (Anzahl der Gleitkommaoperationen pro Sekunde) zeigt sich, dass im Schnitt nur ca. 10 – 20 % der maximalen Gleitkomma-Performance erreicht werden, da der Compiler keine automatische Vektorisierung des C-Codes durchführt. Eine händische Vektorisierung mit Hilfe von Intrinsics wurde nicht vorgenommen. Im Fall von SSE-Instruktionen (Intel Westmere) würde sich dadurch die Performance etwa verdoppeln, bzw. im Fall von AVX-Instruktionen (Intel IvyBridge, Intel Haswell) maximal eine Vervielfachung erfahren. Die Verwendung von einfach genauen Gleitkommazahlen an Stelle von doppelt genauen würde den Löser ebenfalls um einen weiteren Faktor zwei beschleunigen.

Aktuell ist jedoch die erreichte Laufzeit der Lattice-Boltzmann-Simulation kurz genug, und stellt für den Gesamtlauf der 2D-Topologieoptimierung keinen Flaschenhals dar. Relevant werden die genannten weiteren Optimierungsmöglichkeiten jedoch bei der angedachten 3D-Weiterentwicklung des Projekts. Hier wird sich sowohl die Anzahl der Gleitkommaoperationen als auch das Datenaufkommen mindestens verdoppeln.

T. Guess, T. Heidig, B. N. Vu, F. Wein, M. Wittmann, T. Zeiser

1 PetaFlop/s Sustained Performance

Im Herbst und Winter 2013/14 wurden in Kooperation mit dem Leibniz Rechenzentrum mehrere Benchmark-Läufe mit SeisSol auf SuperMUC durchgeführt – zunächst mit dem Ziel einer möglichst guten Performance und Skalierbarkeit auf allen (147.456) verfügbaren Rechenkernen. Bei einem Weak-Scaling-Test mit 20.000 Gitterzellen pro Rechenkern wurde damit bereits im Oktober eine Performance von ca. 980 TeraFlop/s erreicht, womit die „Jagd nach dem PetaFlop“ endgültig eingeläutet war. In einer Extreme

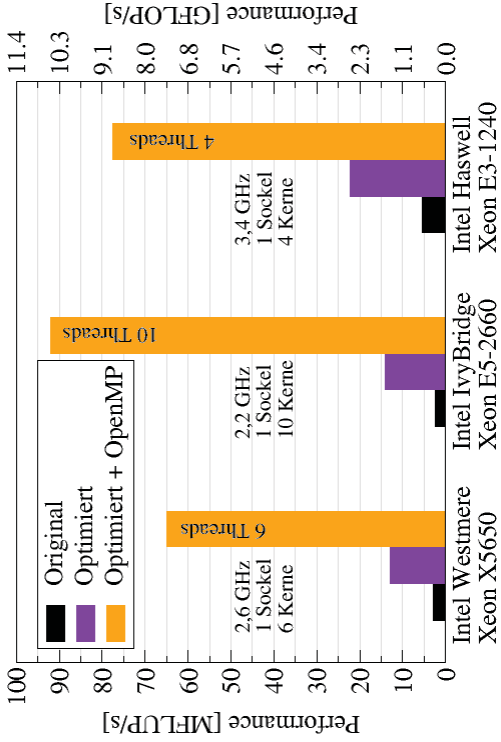


Abbildung 4: Performance des 2D-Lattice-Boltzmann-Strömungslösers für ein Gebiet von 100×100 Knoten vor und nach der Code-Optimierung auf verschiedenen Prozessoren.

rungen der Geometrie zu ziehen. Dazu sind viele Iterationen des Optimierungsalgorithmus nötig, wobei die Strömungssimulation einen signifikanten Anteil der Gesamtrechnenzeit einnimmt.

Durch Zusammenarbeit mit der HPC-Gruppe des RRZE konnte für den existierenden 2D-Lattice-Boltzmann-Strömungslöser eine deutliche Performance-Steigerung erreicht werden: Abhängig vom verwendeten Prozessor ist die optimierte Code-Version nun zwischen 15 und 40 mal schneller als der Original-Code (siehe Abb. 4).

Der wichtigste Punkt der Code-Optimierung war die Reduktion der Datentransfers in der Speicherhierarchie. Dies wurde durch die Verwendung von zwei alternierenden anstatt einem einzigen Gitter sowie die Fusion ein-

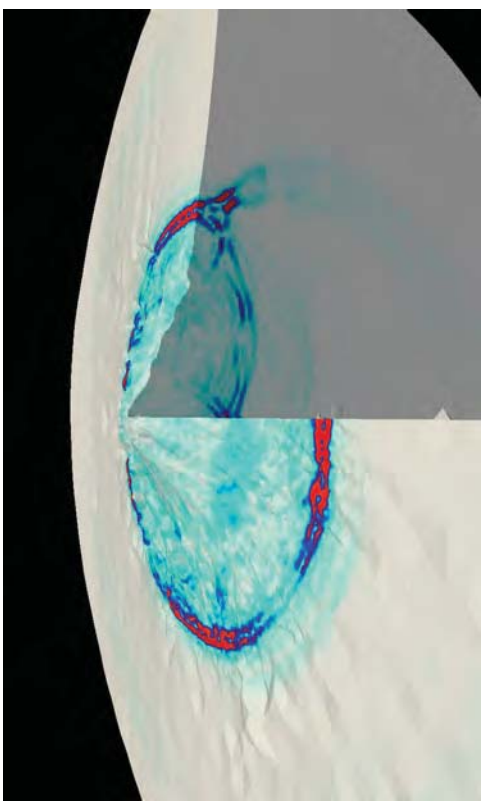


Abbildung 2: Momentaufnahme des seismischen Wellenfeldes am Vulkan Merapi: Reflexionen und Refraktionen an der Topographie und internen Schichten erhöhen dessen Komplexität.

Scaling Period (Ende Januar 2014) wurden die Weak-Scaling-Tests mit einer nochmals in Details verbesserten Version von SeisSol auf nun 60.000 Gitterzellen pro Rechenkern wiederholt. SeisSol erzielte dabei ca. 1.4 PetaFlop/s, wobei ein Problem mit knapp 9 Milliarden Gitterzellen und über 4 Billionen Freiheitsgraden berechnet wurde. Höhepunkt der Tests war jedoch eine Simulation unter Produktionsbedingungen, bei der die Ausbreitung seismischer Wellen im topographisch faszinierenden Vulkan (Merapi) berechnet

wurde (Abbildung 2).

Auf einem Gitter mit knapp 100 Millionen Gitterzellen, wurden dabei in gut drei Stunden Rechenzeit die ersten gut 5 Sekunden dieses Prozesses simuliert. Dass dabei ebenfalls die Marke von 1 PetaFlop/s überschritten wurde, ist als besonderer Erfolg zu werten – Ähnliche Ergebnisse waren bisher nur von Simulationen auf doppelt so großen Maschinen bekannt. Das Spanferkel mit Bier zur entsprechenden Feier wurde schon ratzbutzt aufgegessen!

Danksagung

Wir bedanken uns bei allen, die an der Weiterentwicklung von SeisSol mitgewirkt haben, insbesondere Alexander Breuer, Alexander Heinecke und Sebastian Rettenberger (TU München), Dr. Alice-Agnes Gabriel und Stefan Wenk (LMU München) sowie Dr. Gilbert Brietzke und Dr. Nikolay Hammer (LRZ) und Mikhail Smelyanskiy (Intel Labs). Dem Urheber von SeisSol und Initiator des ASCETE-Projekts, Dr. Martin Käser, möchten wir an dieser Stelle ganz besonders herzlich danken. Ein besonderer Dank geht außerdem an alle Mitarbeiter des LRZ, die unser bei der Durchführung der Läufe auf dem SuperMUC unterstützt haben.

M. Bader, C. Pelties

Performance-Optimierung des Lattice-Boltzmann-Lösers im Verbundprojekt OptiLBM

Im Rahmen des interdisziplinären FAU-Projekts „OptiLBM“ soll ein Katalysatorträger für chemische Reaktoren durch Topologieoptimierung gezielt entworfen werden. Die Strömungssimulation mit Lattice-Boltzmann-Verfahren ist dabei ein zentraler Bestandteil. Durch Code-Optimierungen konnte die benötigte Rechenzeit deutlich reduziert werden.

Ziel des interdisziplinären Projekts „OptiLBM“ der Friedrich-Alexander-Universität Erlangen-Nürnberg (FAU) ist es, einen Katalysatorträger zur Verwendung in chemischen Reaktoren mittels Topologieoptimierung gezielt zu entwerfen. Dazu sind die Kompetenzen mehrerer Disziplinen erforderlich:

Die Arbeitsgruppe von Prof. Stingl (Lehrstuhl für Angewandte Mathematik 2 – AM2; Prof. Leugering) steuert die mathematische Theorie bei während das technische Verständnis zum Bau sowie zur Anwendung von Reaktoren von der Arbeitsgruppe von Prof. Freund (Lehrstuhl für Chemische Reaktionstechnik; Prof. Wasserscheid) eingebracht wird. Die notwendige Performance-Optimierung der Simulation wird von der HPC-Gruppe des RRZE (Prof. Wellein) durchgeführt.



Abbildung 3: Kompetenzen mehrerer Disziplinen

Ausgangspunkt für die Topologieoptimierung der Katalysatorstruktur ist ein homogenes Strömungsgebiet mit festen Ein- und Ausstrombereichen. Das Innere des Gebiets ist nun so zu gestalten, dass eine Zielgröße optimiert wird, z. B. dass der Druckverlust zwischen Ein- und Ausstrom bei gegebenen Strömungsbedingungen minimal wird. Zusätzliche Bedingungen, beispielsweise der maximale Fluidanteil des Gesamtgebiets, können dabei den Lösungsraum einschränken.

Zur Suche nach einer optimalen Struktur mittels Topologieoptimierung wurde im Rahmen der Diplomarbeit von Thomas Guess (AM2) ein iterativer Algorithmus erprobt und in einer ersten Version implementiert. Ein Lattice-Boltzmann-Löser mit geeignetem Porositätsmodell berechnet das Strömungsfeld in der aktuellen Katalysatorstruktur. Ist die Strömungsberechnung konvergiert, wird das Strömungsfeld zusammen mit den Rohdaten der Lattice-Boltzmann-Simulation (d. h. den Particle Distribution Functions) exportiert und darauf aufbauend Druckverlust sowie Gradienten analysiert.

Anhand der Analyseergebnisse wird dann das aktuelle Gebiet angepasst („optimiert“) und die Strömungssimulation erneut gestartet. Diese Iterationsschritt wird so lange wiederholt, bis die Gesamtlösung der Optimierung konvergiert. Eine wichtige Herausforderung besteht darin, aus der Analyse des Strömungsfelds die richtigen Rückschlüsse auf die notwendigen Änderungen