

from 2,616 km² to 1,633 km² and 722 km², respectively. Error ellipse coverage (percentage of GT locations within 90% error ellipses) was 89% without using SSSCs, 91% using model-based SSSCs, and 92% using kriged SSSCs. These results were obtained for source-station paths that sample diverse geological provinces throughout much of Asia. The SSSCs are expected to perform, on average, as well as the test results using the model-based SSSCs, and substantially better for areas where GT calibration data were utilized to refine the SSSCs. Our overall method (of first obtaining model travel times and then kriging with empirical data) is well suited to calibration of any stations in East Asia that have a significant archive of reliably measured arrival times. The method potentially has global application, to greatly improve global bulletins of seismicity without having to rely wholly upon a 3D model of Earth's seismic velocity. But an international long-term effort is needed to build the necessary reference events.

S22F-04 1645h

Relative Earthquake Location Techniques: Tests using both Synthetic and Real Data

Guoqing Lin¹ (gulin@ucsd.edu)

Peter Shearer¹ (shearer@igpp.ucsd.edu)

¹IGPP, SIO, UCSD, 9500 Gilman Dr., La Jolla, CA 92093-0225, United States

Over the years, researchers have developed various techniques for reducing relative errors in earthquake location among nearby events. These techniques rely on the assumption that if the hypocentral separation between two earthquakes is small, travel time differences at a particular station are not biased by the effects of heterogeneity. Here, we compare three different earthquake location techniques: (1) the hypocentral decomposition method of Jordan and Sverdrup (1981), (2) the source-specific station term (SSST) method of Richards-Dinger and Shearer (2000), and (3) the modified double difference method (DD) of Waldhauser and Ellsworth (2001). In principle, all these methods should give approximately the same result for a single compact cluster of events, with significant differences appearing only for more distributed seismicity. We test these methods with both a synthetic data set and actual phase picks and waveform cross-correlation data from the M=5.4 Big Bear California aftershock sequence of February 2003. For the synthetic data, we generate a set of quake locations, station locations and arrival time picks in a simple half-space velocity model. We add random time noise to the data by including contributions of varying size from: (1) normally distributed random picking errors, (2) normally distributed station terms, which are constants for all events recorded by each station, and (3) spatially varying station terms computed by summing travel time anomalies along rays in random 3-D velocity models. In addition, each event is recorded by a random subset of the total set of stations. We quantify the performance of each method by characterizing both the absolute and relative errors in the computed locations as compared to the true locations, experimenting with both compact clusters and more distributed seismicity. The algorithms are also tested on real data, and their performances evaluated on a set of 785 events recorded by the Southern California Seismic Network (SCSN) from a region around the 2003 M=5.4 Big Bear earthquake, for which both phase pick and differential times resulting from waveform cross-correlation are available.

S22F-05 1700h

Location of Earthquakes in Three-dimensional Media Using Repeated, Multiple-Event Locations Over Spatial Grids of Control Points

Gary L. Pavlis¹ (812-855-5141; pavlis@indiana.edu)

Frank L. Vernon² (858-534-5537; vernon@epicenter.ucsd.edu)

¹Indiana University, Department of Geological Sciences 1001 East 10th Street, Bloomington, IN 47405

²University of California, San Diego, Scripps Institute of Oceanography 9500 Gilman Drive, LaJolla, CA 92093-0225

We describe a new approach to the large-scale relocation of seismic catalogs based on a variation of conventional multiple event location methods. We associate every event in a catalog with one or more points in a 3D grid of control points. Events associated with each point are located together with an implementation of the Progressive Multiple Event Location (PMEL) algorithm. PMEL estimates a set of path corrections from a given control point to each seismic station along with a set of revised locations. In the estimation we use a set of matrix projectors that allow the unresolvable (absolute) location bias to be extracted from a 3D reference model and project only components of information the

available data constrain. This allows us to construct empirically derived travel time correction fields in 3D that are highly data adaptive; the density of control points can be very high where seismicity is intense and low where earthquakes are rare. The 3D model is used only to fill in gaps between control points and correct smoothly varying, bias problems. The algorithm is highly parallel and we found major improvements in performance were possible on massively parallel computers. We applied this approach to data from the Anza network and the Tien Shan (Ghengis) experiment catalog. With the Anza data we achieved remarkable improvements in data fit with most relocation residuals smaller than 0.02 s which is approximately 1/10 of the catalog residuals. In contrast the variance reduction in the Tien Shan data was much smaller (50%) with 1 s scale residuals common. This difference is attributed to measurement precision. The Anza catalog is dominated by impulse, local P and S phases while the Tien Shan catalog is dominated by emergent, regional phases.

S22F-06 1715h

Use of 3-D Velocity Models and Ray Tracing in the Double Difference Earthquake Location Algorithms: Application to a Data set from the Mygdonia Basin (Northern Greece)

Odysseus C. Galanis¹ (30-231-0-99-1406; ogalanis@lemnios.geo.auth.gr)

Constantinos B. Papazachos¹ (30-231-0-99-8510; costas@lemnios.geo.auth.gr)

Panagiotis M. Hatzidimitriou¹ (30-231-0-99-8505; takis@lemnios.geo.auth.gr)

Emmanuel M. Scordilis¹ (30-2310-0-99-1411; manolis@lemnios.geo.auth.gr)

¹Geophysical Laboratory, School of Geology, Aristotle University of Thessaloniki, P.O. Box 352-1, Thessaloniki 54006, Greece

In the past years there has been a growing demand for precise earthquake locations for seismicity, tectonics and seismic hazard studies. Recently this has become possible because of the development of sophisticated location algorithms, as well as hardware resources. This is expected to lead to a better insight of seismicity in the near future. A well-known technique, which has been implemented as a computer program and has been widely used for relocating earthquake data sets, is the double difference algorithm. In its original implementation it makes use of a one-dimensional ray tracing routine to calculate seismic wave travel times and partial temporal derivatives of source parameters. We have modified the implementation of the algorithm by incorporating a three-dimensional velocity model and ray tracing in order to locate a set of earthquakes in the area of the Mygdonia Basin (Northern Greece). This area is covered by a permanent regional network and occasionally by temporary local networks. The velocity structure is very well known, as the Mygdonia Basin has been used as an international test site for seismological studies since 1993, which makes it an appropriate location for evaluating earthquake location algorithms, with the quality of the results depending only on the quality of the data and the algorithm itself. The new earthquake locations reveal details of the area's structure, which are blurred, if not misleading, when resolved by standard (routine) location algorithms.

S22F-07 1730h

Fine-Scale Structure of the San Andreas Fault and Location of the SAFOD Target Earthquakes

C. Thurber¹ (thurber@geology.wisc.edu)

H. Zhang¹ (hjzhang@geology.wisc.edu)

S. Roecker² (roecker@rpi.edu)

W. Ellsworth³ (ellsworth@usgs.gov)

P. Malin⁴ (malin@duke.edu)

¹UW-Madison, 1215 W. Dayton St., Madison, WI 53706, United States

²RPI, 110 8th Street, Troy, NY 12180, United States

³US Geological Survey, 345 Middlefield Rod, Menlo Park, CA 94025, United States

⁴Duke University, Box 90235, Durham, NC 27708, United States

Arrival-time data from about 900 local earthquakes and several dozen shots are inverted for earthquake locations and three-dimensional Vp and Vp/Vs structure of the San Andreas fault (SAF) zone at Parkfield, CA, using several different tomography methods. Data are included from a temporary array of surface seismic stations installed around the SAFOD site, a vertical array of geophones installed in the 2-km-deep SAFOD

Pilot Hole drilled in summer 2002, and permanent network stations. The inversion methods applied are conventional ray-tracing tomography, conventional finite-difference tomography, and double-difference tomography. For double-difference tomography, we use both a regular grid and an adaptive grid based on tetrahedra and Voronoi diagrams. The primary features of the different velocity models are generally consistent with each other. A low-Vp zone (about 25% slower) adjacent to the fault trace penetrates to a depth of as much as 10 km. Nearly all the earthquakes occur almost directly beneath the surface trace of the SAF, on the edge of a zone of high horizontal Vp gradient. Regions of anomalously high Vp/Vs (> 2.0) are found at shallow depths along and northeast of the SAF. These high-Vp/Vs anomalies can be associated with features in an existing resistivity model that are interpreted to represent fluid-saturated zones. We locate several magnitude 1 to 2 earthquakes at depths < 3 km and locations that are on the order of 100 meters southwest of the surface fault trace that are potential target events for drilling of the main SAFOD hole. We discuss a range of tests of our absolute and relative location capability for events in the target region, and identify additional steps needed to reduce location uncertainties to < 100 m and ultimately < 10 m.

S22F-08 1745h

High-Resolution Subducting Slab Structure Beneath Northern Honshu, Japan, Revealed by Double-Difference Tomography

Haijiang Zhang¹ (hjzhang@geology.wisc.edu);

Clifford Thurber¹ (thurber@geology.wisc.edu);

David Shelly² (dshelly@pangea.Stanford.EDU);

Satoshi Ide³ (ide@eps.s.u-tokyo.ac.jp); Gregory

C. Beroza² (beroza@pangea.Stanford.EDU); Akira

Hasegawa⁴ (hasegawa@aob.geophys.tohoku.ac.jp)

¹University of Wisconsin-Madison, Department of Geology and Geophysics, 1215 W. Dayton St., Madison, WI 53706, United States

²Stanford University, Department of Geophysics, 397 Panama Mall, Stanford, CA 94305, United States

³University of Tokyo, Department of Earth and Planetary Science, Tokyo 113-0033, Japan

⁴Tohoku University, Research Center for Prediction of Earthquakes & Volcanic Eruption, Graduate School of Science, Sendai 980-8578, Japan

Double seismic zones separated by 30 km exist at intermediate depths of 50-160 km in northeast Japan. The upper and lower plane of earthquakes occur mainly in oceanic crust and mantle, respectively, and are assumed to be induced by dehydration reactions of various hydrous minerals. We obtain the seismic structure of the subducting slab with unprecedented resolution beneath Northern Honshu, Japan by applying newly developed double-difference tomography to the two planes of seismicity of the double seismic zone, which provides evidence supporting the above hypothesis. The upper plane consists of two distinct regions: one with relatively high Vp/Vs ratio (1.77-1.8) at 60-90 km depth and one with low Vp/Vs ratio (1.7-1.77) at 90-110 km depth. These two regions may correspond to the transformations of metabasalt/metagabbro to blueschist and blueschist to eclogite, respectively. There is a conspicuous anomaly of very low Vp/Vs ratio (1.6-1.7) associated with the lower plane, in sharp contrast with high Vp/Vs ratio (1.8-1.85) in the region between the two planes. The forsterite-enstatite-H₂O system from serpentine dehydration for the lower plane and possible partial hydration of the region between the two planes of earthquakes may explain the observed features.

S22G MCC: 2007 Tuesday 1600h

Rise of the Machines: Wave Propagation Theory

Presiding: C J Thomson, Queen's University; K T Nihei, Lawrence Berkeley National Laboratory

S22G-01 1600h

Coherent-state analysis of the seismic head-wave problem: an overcomplete representation and its relationship to rays and beams

Colin J. Thomson ((613) 533 6178; thomson@geol.queensu.ca)

Queen's University, Geological Sciences & Geological Engineering, Kingston, ON K7L 3N6, Canada

The coherent-state transform (CST) is a Gaussian-windowed Fourier transform and compared to the usual plane-wave expansion of seismic wavefields it provides an overcomplete basis of partial waves. These overcomplete partial waves have associated rays and the relationship of these rays to those of a physical wavefront is explained here by appealing to a familiar analytic example. The exact CST of the standard Fourier plane-wave summation for a point-source primary wave can be interpreted as a sum of damped plane waves forming a focussed or beam-like wavefield. This primary-wave CST can also be approximated as a bundle of complex rays in a position-slowness space with higher dimension than that usually considered, a consequence of the overcompleteness. The higher-dimensional ray spreading controls the coherent-state (CS) amplitude. This complex-ray bundle in turn can be approximated as a paraxial Gaussian beam carried by a real ray associated with the 'physical' wavefront. The exact CST of the standard plane-wave summation for a point-source reflection shows how the complex-ray and paraxial-beam approximations generalize to interfaces. Around a critical angle the standard plane-wave summation has an asymptotic form involving the Weber function and this function also necessarily arises in the complex-ray and paraxial-beam approximations for the reflection CST. The inverse CST giving the point-source wavefield is a sum of coherent states and those contributing rays or paraxial beams which are incident around the critical angle should be given a reflection coefficient involving the Weber function. CS beams which are incident away from the critical angle collect a peripheral branch-point diffraction. In general, both types of critical-ray/branch-point signal should be included if the integrated CS response is to correctly describe the head wave. An advantage of the CST is that it combines with ray theory for gradually-varying media to give a convenient solution to the caustic and pseudocoustic problems. The general idea of such overlapping partial-wave expansions, their flexibility and the smoothing they impart may have other applications in seismic analysis and processing.

S22G-02 1615h

The Physical Basis of Lg Generation by Explosion Sources

Jeffrey L Stevens¹ ((858) 826-1635;

Jeffrey.L.Stevens@saic.com); G. Eli Baker¹ ((858) 826-1606; glenn.e.baker@saic.com); Heming Xu¹ ((858) 826-1641; heming.xu@saic.com); Theron J. Bennett¹ ((703) 247-4156; theron.j.bennett@saic.com); Norton Rimer¹ ((858) 826-1654; norton.rimer@saic.com); Steven M. Day² ((619) 594-2663; day@moho.sdsu.edu)

¹Science Applications International Corporation, 10260 Campus Point Drive, San Diego, CA 92121, United States

²San Diego State University, Dept. Geological Sciences MC-1020 5500 Campanile Dr., San Diego, CA 92182, United States

We use observations of explosion-generated Lg together with numerical modeling to determine how underground nuclear explosions generate shear wave phases. This question is fundamental to how Lg phases are interpreted for use in explosion yield estimation and earthquake/explosion discrimination. We analyze several explosion data sets: 1) Degelen Mountain explosions recorded at distances less than 100 km and corresponding recordings at Borovoye (BOR) at a distance of approximately 650 km; 2) recordings from Russian deep seismic sounding experiments; 3) Nevada Test Site (NTS) explosion sources including the Nonproliferation Experiment (NPE) and nuclear tests covering a range of source depths and media properties. Observations that constrain possible models include: 1) Small Sg phases are observed at distances from less than a km to 10s of km; 2) a strong Rg phase can persist to 100s of km; 3) at 10s to 100s of km, Sg or Lg remain distinct from Rg, with no indication of scattered energy preceding Rg; 4) Sn is impulsive and large for Degelen explosions recorded at BOR; 5) Shear wave phases are larger on horizontal than vertical seismograms. We model the overburied NPE, and an underburied Degelen explosion, using point sources and two-dimensional nonlinear finite difference calculations. A small Sg distinct from large Rg near the Degelen explosion and large Lg and Sn at regional distance are consistent with the signals from a shallow, axisymmetric CLVD source. This source has an S node in the horizontal direction, and therefore makes only a small direct S wave at the near regional stations, however it is a strong generator of S at takeoff angles corresponding to the crustal phases Lg and Sn. Large Lg from the NPE can be generated from just a point explosion due to the low velocity of the source medium. We are considering three candidate mechanisms for explosion-generated Lg: 1) Direct generation by the explosion source, where the explosion is modeled as a point compressional source; 2) Secondary generation by the explosion source, where Lg is generated primarily by the nonspherical parts of the explosion source, with strong influence from the free surface; and 3) Rg scattering. The observations discussed above favor explanation number 2 in a high velocity source medium and a combination of 1 and 2 in a low velocity

medium. The spherically symmetric part of an explosion source in a high velocity medium generates very little Lg, and therefore does not explain the observations. However, the observed vertical and radial shear waves can be explained by adding a CLVD, which is the lowest order non-spherical correction to the spherical source. That Rg persists to such large distances argues against explanation 3. More difficult to explain is the observed presence of Lg on the tangential components. This cannot be modeled with an axisymmetric source, and must be due to either a non-axisymmetric source effect, or to polarization changes along the source to receiver path.

S22G-03 1630h

Regional and Upper-Mantle Structure Model of Peninsular India and Relevance to High-Frequency Lg Seismograms

Chandan K. Saikia (626-449-7650; chandan.saikia@urscorp.com)

URS Group, Inc., 566 El Dorado Street, Pasadena, CA 91101, United States

The primary objective of this study is to develop regional and upper-mantle crustal structures of Peninsular India for estimating source parameters of smaller seismic events. To this end, we modeled regional seismograms generated by the May 21, 1997 Jabalpur earthquake in central India which occurred in the lower crust at a depth of 35 km. It was recorded at many teleseismic stations with usable signal-to-noise ratio, including several broadband regional stations in India. Teleseismic P waves were used to establish the complexity of the source process, focal mechanism and source depth of the earthquake by modeling amplitude and travel times of the pP and sP depth phases relative to the P wave onsets. Regional waveforms from this master event with known source parameters were used to develop corresponding Greens' functions for the paths to stations, namely Bhopal (BHPL), Bilaspur (BLSP) and Hyderabad (HYB). The complexity in teleseismic P and sP can also be observed in Pn and sPn with excellent compatibility. These same features can be identified in upper-mantle triplications at LSA ($\Delta=12^\circ$), NIL ($\Delta=12^\circ$), CHTO ($\Delta=18^\circ$), and AAK ($\Delta=20^\circ$) when modeled with existing shield models. The upper-mantle Pn phases (i.e., pP, sP, sPn, sSP, sPP, sPPP etc.) can be positively identified in these observations and were found to be sensitive to the presence of gradient in the transition zones, especially beneath the 400 km discontinuity. The usefulness of these Green's functions is demonstrated by modeling the regional and sparse teleseismic data for a small earthquake (April 4, 1995) located near the Indian Nuclear Test Site at Pokhran. This study also establishes a regional waveguide for the path from the Pokhran Test Site to station Nilore (NIL) in Pakistan. We have also extended these regional models to have the capability of generating high-frequency Lg waves by introducing thin layers in the crust with alternating high and low-velocity distribution.

S22G-04 1645h

Accurate Multi-Phase Traveltimes in 2-D Layered Media Using a Fast Marching Scheme With Source Grid Refinement

Nicholas Rawlinson¹ (61-2-6125-0339; nick@rsees.anu.edu.au)

Malcolm Sambridge¹ (61-2-6125-4557; malcolm@rsees.anu.edu.au)

¹Research School of Earth Sciences, Australian National University, Canberra, ACT 0200, Australia

The accurate prediction of seismic traveltimes in layered media is required in many areas of seismology. In addition to simple refractions and reflections, complex phases comprising numerous transmission and reflection branches may exist; for instance, the so-called "multiples" frequently identified in marine reflection seismology. We present a grid-based method for the accurate determination of multi-phase traveltimes in layered media of significant complexity. A finite difference eikonal solver known as the Fast Marching Method (FMM) is used to track wavefronts within a layer. FMM is a fast and unconditionally stable upwind scheme that is well suited to complex models, and can be used sequentially to track the multiple refraction and/or reflection branches of virtually any required phase. Although FMM was initially introduced as a first-order scheme, higher order operators can be used. A mixed-order scheme that preferentially uses second-order operators, but reverts to first-order operators when the required upwind traveltimes are unavailable, is one possibility. Despite improved accuracy, this scheme still suffers from first-order convergence due to high wavefront curvature and first-order accuracy in the vicinity of the source. To overcome this problem, we implement local grid refinement about the source. In order to retain stability, the edge of the refined grid conforms

to the shape of the wavefront, so that information only flows out of the refined grid, and never back into it. Application of our new scheme to complex velocity media shows that grid refinement typically improves accuracy by an order of magnitude, with only a small increase in computation time (~5%). Significantly, first-order convergence is replaced by near second-order convergence, even in media with velocity contrasts as large as 8:1. In one example, with a velocity grid defined by 257,121 nodes, reflection traveltimes from a strongly undulating interface were calculated with an error of only 0.001% in approximately 5 s of CPU time (on a Sun Blade 150). This level of accuracy is sufficient for calculating meaningful amplitude values via solution of the transport equation.

S22G-05 1700h

Absorbing Boundaries in Frequency-Space and Wavenumber-Time Domains

Francisco J Sanchez-Sesma¹ (sesma@servidor.unam.mx)

Carlos I. Villa-Velazquez¹

¹Instituto de Ingenieria, UNAM, Cd. Universitaria, Circuito Escolar s/n, Coyoacan, DF 04510, Mexico

Absorbing boundary conditions are used in a variety of disciplines to simulate the radiation of energy to the exterior of a numerically-studied interior domain. The scattering of sound, elastic or electromagnetic waves by arbitrary shaped objects may require terminating the interior domain with adequate boundary conditions. In this work a perfect absorbing boundary condition for 2D elastic wave propagation is formulated in the frequency-wavenumber spectral domain (f-k). Using Fourier transformations, the corresponding results are obtained in both the frequency-space (f-x) and the wavenumber-time domains, respectively. The singular results in the space-time (x-t) domain are discussed as well. The absorbing boundary condition can be regarded as an interface operator that gives the normal derivative in terms of boundary values. Thus it is called Dirichlet to Neumann (DtN) operator, a term that has been coined in the mathematical community. Within the framework of elasticity and from physical considerations, this condition can be expressed in terms of tractions and displacements and the DtN operator is simply the stiffness of the exterior region. This term comes from engineering usage. Most of the results presented here correspond to the 2D antiplane SH case but the ideas discussed have their vector counterparts for the in-plane P and SV waves and the 3D problems. It is found that our expressions in the space-time (x-t) domain appear naturally in the classical dynamic rupture problem. Partial support is acknowledged to DGAPA-UNAM under grant IN121202 and CONACYT-Mexico under project NC204.

S22G-06 1715h

Amplitudes and Traveltimes of Waves in Random Media

Adam M Baig¹ (609 258 1504; abaig@princeton.edu)

F. A. Dahlen¹ (fad@princeton.edu)

¹Department of Geosciences, Princeton University, Princeton, NJ 08544, United States

We measure traveltimes and amplitudes from 'ground-truth' seismograms, computed using a numerical wave propagation code to compare with finite-frequency and ray-theoretical values for these quantities. Ray-theoretical traveltimes become invalid whenever the scale-length of 3-D heterogeneity is smaller than half the maximum width of the Fresnel zone; in contrast, ray-theoretical amplitudes have a much more restricted range of validity, that the scale length should be less than one Fresnel zone width. Finite-frequency theory gives better results for amplitudes, suffering no observable degradation for small-scale media for the weakest heterogeneity considered, but suffering appreciable misfit in more heterogeneous media. Using these finite-frequency expressions for traveltime, we derive expressions for the expected variances of traveltimes and amplitudes that act, in most cases, as extensions to the ray-theoretical expressions. Finally, we propose using the amplitude variance as a criterion for delineating the validity of these linear approximations. For traveltimes, provided one rejects waveforms that do not provide a good cross-correlation traveltime, the remaining data are linearly related to the model over the values of theoretical amplitude variance that we probe in this experiment. Amplitudes do not behave as well: when the theoretical amplitude variance rises above 0.1, significant nonlinearities start to invalidate our linear approximation.

S22G-07 1730h

An Efficient Approach for Computing the Frequency Response of Seismic Waves in Heterogeneous, Anisotropic Viscoelastic Media With FDTD+PSD Modeling

Kurt T. Nihei¹ (510-486-5349; ktnihei@lbl.gov)Seiji Nakagawa¹ (510-486-7894; snakagawa@lbl.gov)¹Lawrence Berkeley National Laboratory, 1 Cyclotron Rd., MS 90-1116, Berkeley, CA 94720, United States

Computation of the frequency response (phase and magnitude) of seismic waves propagating in heterogeneous, anisotropic, viscoelastic media is required for a number of seismology efforts, including frequency-domain full-waveform inversion and basin modeling. When the frequency response is required at a limited number of locations, it can be computed efficiently with a memory variable, staggered grid finite difference time domain (MV-SG-FDTD) code by taking the Fourier transform of the received waveforms generated by a broadband source. However, when the frequency response is required at many or all grid locations, as in frequency-domain full-waveform inversion, the memory requirements for storing the waveforms make this approach prohibitive. An alternative approach is to compute the frequency response by formulating the SG-FDTD equations in the frequency domain. For O(4) accuracy spatial differencing, the resulting 2-D system of implicit equations for the unknown particle velocities and stresses at the finite difference cell locations yields a complex system matrix of order equal to $[5 \times M \times N]^2$, where M and N are the number of finite difference cells in the x and z directions, respectively, and the 5 comes from the five unknowns in a 2-D staggered grid finite difference stencil (2 particle velocities, and 3 stresses). The complex matrix equation has the following properties: (1) sparse (total number of unknowns equals $[45 \times M \times N - 20 \times M - 20 \times N]$), (2) banded (4 bands of sparse 5×5 submatrices on each side of the main diagonal), and (3) non-Hermitian. We found that the resulting system matrix is typically much too large for direct matrix solvers such as sparse LUD for seismic problems of reasonable size. Additionally, our attempts to solve the frequency domain system of implicit equations using iterative sparse matrix solvers, such as QMR (quasi minimum residual method), tended to suffer from poor convergence problems when using standard preconditioning (e.g., Jacobi). In this presentation, we demonstrate that an efficient approach for computing the frequency response of a heterogeneous, anisotropic, viscoelastic medium can be achieved by running an explicit MV-SG-FDTD code with a harmonic wave source out to steady-state, and then extracting the magnitude and phase from the transient data via a phase sensitive detection algorithm (PSD). The PSD algorithm requires integration over only several cycles of the waveform to obtain accurate phase and magnitude estimates. Because this integration is performed on-the-fly, there is no need to store waveforms at the grid locations. Preliminary tests indicate that it should be possible to superimpose multiple frequencies at the source and extract the magnitude and phase for each frequency using the PSD algorithm, opening up the potential for obtaining the multi-frequency response of a heterogeneous, anisotropic, viscoelastic medium with a single FDTD run.

S22G-08 1745h

Numerical Simulation of Time-Dependent Wave Propagation Using Nonreflective Boundary Conditions

Dorin Ionescu¹ (+61 7 3365 3464; ionescu@quakes.uq.edu.au)Hans Muehlhaus¹ (+61 7 3365 4783; Muehlhaus@quakes.uq.edu.au)¹Earth Systems Science Computational Centre and The Australian Computational Earth Systems Simulator, The University of Queensland St. Lucia Campus, Brisbane, QLD 4072, Australia

Solving numerically the wave equation for modelling wave propagation on an unbounded domain with complex geometry requires a truncation of the domain, to fit the infinite region on a finite computer. Minimizing the amount of spurious reflections requires in many cases the introduction of an artificial boundary and of associated nonreflecting boundary conditions. Here, a question arises, namely which boundary condition guarantees that the solution of the time dependent problem inside the artificial boundary coincides with the solution of the original problem in the infinite region. Recent investigations have shown that the accuracy and performance of numerical algorithms and the interpretation of the results critically depend on the proper treatment of external boundaries. Despite the computational speed of finite difference schemes and the robustness of finite elements in handling complex

geometries the resulting numerical error consists of two independent contributions: the discretization error of the numerical method used and the spurious reflection generated at the artificial boundary. This spurious contribution travels back and substantially degrades the accuracy of the solution everywhere in the computational domain. Unless both error components are reduced systematically, the numerical solution does not converge to the solution of the original problem in the infinite region. In the present study we present and discuss absorbing boundary condition techniques for the time-dependent scalar wave equation in three spatial dimensions. In particular, exact conditions that annihilate wave harmonics on a spherical artificial boundary up to a given order are obtained and subsequently applied in numerical simulations by employing a finite differences implementation.

S31A MCC: 3009 Wednesday 0800h

Stress Transfer, Triggered Earthquakes, and Time-Dependent Seismic Hazard I (joint with G, T, NG)

Presiding: S Steacy, University of Ulster; E S Cochran, Institute of Geophysics and Planetary Physics, University of California, Los Angeles

S31A-01 0800h INVITED

A Fresh Look at the Triggering of Earthquake Pairs, Such as the Landers-Big Bear, Landers-Hector Mine, Izmit-Duzce, and Nenana-Denali, and March-May 1997 Kagoshima Events

Shinji Toda¹ (81-298-61-3743; s-toda@aist.go.jp)Ross S. Stein² (650 329 4840; rstein@usgs.gov)¹Active fault Research Center (AIST), Higashi 1-1, Tsukuba, CA 305-8567, Japan²U.S. Geological Survey, MS 977, Menlo Park, CA 94025, United States

Recent precise hypocentral information provides opportunities to study how secondary earthquakes are spatially and temporally triggered by mainshocks. Most studies regard the second shock in the pair as an aftershock triggered exclusively by the mainshock, with its likelihood decreasing with time after the first event. Here we argue that small aftershocks occurring near the future hypocenter of the second shock play an important role in enhancing the likelihood of the pair. For such multiple-order triggering calculations, we incorporate ate/state friction of Dieterich (1996) into stress transfer. The results may answer two riddles: The first is to ask why there is a net probability gain during the aftershock sequence when there must be equal areas of stress enhancement (trigger zones) and stress depression (shadows) near the mainshock. The Omori law describes only the behavior in the trigger zones. The answer appears to be that the net probability gain occurs because the seismicity rate increases exponentially in response to the stress change, so the seismicity rate gain in the trigger zones dwarfs the seismicity rate drop in the shadows. The second riddle is to ask how small aftershocks triggered by the mainshock affect the likelihood of the second large earthquake in the pair. The answer is that because the state variable (gamma) in the Dieterich earthquake formulation has already plummeted as a result of the mainshock, the effect of small aftershocks on the nearby seismicity rate is highly amplified relative to the effect of the mainshock, even if the mainshock and the small aftershock change the stress by the same amount. The implication is that that probability of the second shock is not necessarily highest immediately after the mainshock, as we typically assume. Instead, it might be highest when the distribution of first-order aftershocks expands (as also predicted by rate/state friction) to a potential earthquake source, such as an active fault. A more complete probability calculation would thus need to include the effects of smaller aftershocks, as advocated by Felzer et al (2002), a very challenging task.

URL: http://quake.wr.usgs.gov/research/deformation/modeling/refs/ross_refs.html**S31A-02 0815h**

Slip Partitioning and the Regional Stress Field

David Bowman^{1,2} (714 278-5436; dbowman@fullerton.edu)Geoffrey King² (king@ipgp.jussieu.fr)Yann Klinger² (klinger@ipgp.jussieu.fr)Paul Tapponnier² (tappon@ipgp.jussieu.fr)¹California State University, Fullerton, 800 N. State College Blvd, Fullerton, CA 92834-6850, United States²Institut de physique du Globe de Paris, 4, place jussieu, Paris 75252, France

Oblique motion along tectonic boundaries is commonly partitioned into slip on several faults with different senses of motion. This partitioning can be explained by the upward propagation of a localized fault at depth. The static stress field ahead of the propagating fault separates into zones of predominantly normal, reverse and strike-slip faulting. Elastic approximations to plastic behavior have been used to explain the distribution of faults observed along the San Andreas (California, USA) and the Haiyuan faults (Tibet, China). Slip partitioning in a single earthquake rupture has also been documented for the 2001 Kokoxili, Tibet earthquake. This process has important implications for regional stress fields in areas where the strain is partitioned onto structures with different senses of motion. In particular, the notion of a partitioned stress field driven by slip at depth can be used to generate a realistic background stress field for models of Coulomb stress interactions. This can be incorporated into more accurate calculations of the stress interactions in aftershock sequences of large earthquakes, which is crucial to determining real-time aftershock probabilities.

S31A-03 0830h

Detecting Seismicity Shadows at Short Time Scales

David Marsan (david.marsan@univ-savoie.fr)

LGIT Universite de Savoie, Campus Scientifique, Le Bourget du Lac 73376, France

Stress changes following the occurrence of an earthquake can lead to significant changes of the seismicity rate. While triggering (eg. aftershocks on the main fault) is clearly observed, regions of seismicity rate decrease can be harder to detect. We here address this detectability issue, which becomes particularly severe at short (eg. days-weeks) time scales. A robust method is proposed in order to estimate the probability that a given region has experienced a seismicity rate decrease. This probability is generally (for large, recent California earthquakes) found to be low: seismicity shadows are significantly not as common as activated regions. The role of crustal heterogeneities in promoting triggering is discussed.

S31A-04 0845h

Testing the stress shadow hypothesis

Karen R. Felzer¹ (felzer@seismology.harvard.edu)Rachel E. Abercrombie² (rea@bu.edu)Emily E. Brodsky¹ (brodsky@ess.ucla.edu)¹Department of Earth and Space Science, Univ. Calif. Los Angeles, 1708 Geology Building, 595 Charles Young Dr. East, Los Angeles, CA 90095-1567, United States²Department of Earth Sciences, Boston University, 685 Commonwealth Ave., Boston, MA 02215

According to current theoretical understanding an earthquake may decrease the rate of other earthquakes in a particular area (create a stress shadow) if earthquakes influence each other via static stresses, but not if they primarily interact via dynamic stresses (Toda and Stein, 2003). Using the California ANSS catalog from 1978 to 2001, we find that the occurrence of stress shadows is difficult to prove, indicating that dynamic earthquake interaction may be a viable alternative. Stress shadows are often identified by measuring seismicity rates before and after a large earthquake and then noting where rate decreases are larger than those expected from a random Poissonian process (Matthews and Reasenberg, 1988). These counts by themselves offer no proof that the mainshock is the cause of the observed rate decreases, however. Marsan (2003) projected the decay rates of pre-existing aftershock sequences to demonstrate that in many cases measured "shadows" were actually caused by ongoing localized aftershock sequence decay. We use an alternative method to verify the results of Marsan (2003) to test whether rate decreases after large earthquakes can be attributed in general to independent causes. We compare the degree and extent of seismicity rate changes around the times of large earthquakes and around other points in time. Like the method of Marsan (2003) this method does not require declustering, often a somewhat subjective process. Preliminary results indicate that while the amount of localized rate increases after large earthquakes is much larger than at other times, the amount of rate decreases may not be significantly different. In Southern California, for example,