

Earthquake Early Warning and the Physics of Earthquake Rupture

By

Gilead Wurman

A dissertation submitted in partial satisfaction of the

requirements for the degree of

Doctor of Philosophy

in

Earth and Planetary Science

in the

Graduate Division

of the

University of California, Berkeley

Committee in charge:

Professor Richard M. Allen, Chair

Professor Douglas S. Dreger

Professor Steven D. Glaser

Spring 2010

Abstract

Earthquake Early Warning and the Physics of Earthquake Rupture

by

Gilead Wurman

Doctor of Philosophy in Earth and Planetary Science

University of California, Berkeley

Professor Richard M. Allen, Chair

One of the great debates in seismology today revolves around the question of whether earthquake ruptures are self-similar, cascading failures, or whether their size is somehow predetermined at the start of the rupture. If earthquakes are self-similar there is theoretically no way to determine the magnitude of an event until the rupture has completely terminated, while if it is deterministic the magnitude should be immediately discernible. Recent advances in Earthquake Early Warning methodologies provide new insight into the fundamental physics of earthquake rupture and highlight the importance of understanding the answer to this question.

Observations of the amplitude and frequency content of early P-wave arrivals suggest that some information about the final size of an earthquake is already present within a few seconds of the initiation of rupture, in agreement with a host of other observations that show a degree of scaling between large and small earthquakes. While this suggests that earthquakes are deterministic, there is likewise a large body of work, both observational and model-based, that indicates that this is not true and earthquakes are self-similar.

This work documents the process of calibrating and testing the ElarmS Earthquake Early Warning methodology in northern California on the Northern California and Berkeley Digital Seismic Networks. In the process the work adds to the body of observations which show a dependency on event magnitude of P-wave frequency content and amplitude. These observations are corroborated with a new set of independent observations of kinematic slip distributions. These new observations indicate that the early slip on a fault also scales with magnitude and suggest again that earthquakes are not entirely self-similar cascading events.

In an effort to assign a physical mechanism to the observations of scaling, both in P-waves and in kinematic slip inversions, a hypothetical model is tested wherein the intensity of the early rupture imparts more or less energy to the rupture front and affects the likelihood of the rupture continuing or dying out in the face of unfavorable conditions further along the fault plane. The results of testing this hypothesis are somewhat equivocal, but they are suggestive of the likely truth, that earthquakes exhibit aspects of both deterministic and cascading rupture to some degree. Understanding the details of the interplay between these two aspects is crucial to the successful application of Earthquake Early Warning systems, especially in rare large earthquakes for which there is little empirical data on the performance of these systems.

Table of Contents

1. Introduction	1
2. Toward Earthquake Early Warning in Northern California	5
2.1. Abstract	6
2.2. Introduction.....	6
2.3. The ElarmS Methodology	8
2.3.1. Calibration dataset	8
2.3.2. Triggering and event location	9
2.3.3. Magnitude from predominant period	9
2.3.4. Magnitude from P-wave peak amplitude.....	11
2.3.5. Data integration and magnitude determination	12
2.3.6. Ground motion prediction.....	15
2.3.7. Simulating ElarmS.....	16
2.4. ElarmS Performance	16
2.4.1. Performance of non-interactive processing	16
2.4.2. Two Bay Area scenario events	19
2.5. Improving ElarmS Performance	23
2.6. Conclusions	24
2.7. Acknowledgments	25
2.8. Figures.....	26
3. Statistical Testing of Theoretical Rupture Models Against Kinematic Inversions	35
3.1. Abstract	36
3.2. Introduction.....	36
3.3. Method	38
3.4. Relationship between early and final moment	39
3.5. Effects of Size Distribution, Bias and Stress Drop	42
3.5.1. Size Distribution	42
3.5.2. Bias.....	42
3.5.3. Stress drop	43
3.6. Variability due to inversion data and methods	45
3.7. Variability within individual events.....	47
3.8. Conclusions	48
3.9. Data and Resources.....	49
3.10. Acknowledgements	49
3.11. Tables and Figures.....	50
4. Modeling the Effect of Early Rupture on Earthquake Magnitude	59
4.1. Abstract	60
4.2. Introduction.....	60
4.3. Method	62

4.4. Modeling in 2D	66
4.5. Modeling in 3D	67
4.6. Comparison of Parametric Values	69
4.6.1. Regularity	70
4.6.2. Constraint	70
4.6.3. Correctness	71
4.7. Conclusions	72
4.8. Acknowledgements	73
4.9. Tables and Figures.....	74
5. Conclusion.....	80
6. References	83

Acknowledgements

My advisor, Richard Allen, deserves a great deal of credit for having the patience and tenacity to continue advising me these last five years. I am not an easy grad student to direct and collaborate with, and it is to his credit that we are ending this endeavor on good terms. My thanks also to Steve Glaser, my outside committee member and to Doug Dreger, my final committee member and my M.S. advisor in a previous graduate life. I think of all the faculty in the department Doug has been my role model and I hope we continue our collaboration, and indeed our friendship into the future. I also wish to thank David Oglesby, my collaborator at UC Riverside for his great help in writing the bulk of this dissertation, and for encouraging words at times when I thought my work would lead nowhere.

I think most, if not all, grad students experience a period of doubt sometime during their time in grad school. I don't know if I experienced more than the average grad student, but I certainly had my fair share of "why am I still doing this?" moments. In those times I relied on two things to keep me in the game: my friends and family, and my own stubbornness. I suppose I have my father to thank for the stubbornness, and my mother for moderating it enough so that I could have such wonderful friends during my time at Berkeley. And indeed there are many such friends. Within the department, Alyssa provided me with moral support, good political debates and a ready couch so that when I didn't feel like working I could do so in her office. Trey's political views and mine are too similar to have debates, but likewise he was always ready for lengthy off-topic debates and Trey more than anyone gave me hope that I could indeed hold down a job, care for a young child, and finish my dissertation all at the same time. Outside the department my flight instructor, Mal, at one point told me "if you quit out of your Ph.D. I'll never fly with you again." For me, that's motivation in its highest form. To George and Mike, my superiors at Seismic Warning Systems, I offer my thanks for their patience and indulgence in allowing me to finish my dissertation while working for them full-time. Finally I have to express my deep love and gratitude to my wife, Lori, who supported me both morally and financially through the past five years and served as a role model of academic excellence. It is thanks to her more than anyone else that this dissertation every saw the light of day.

1. Introduction

Earthquakes are among the most powerful natural phenomena on Earth, and when large earthquakes have coincided with human populations they have caused some of the deadliest natural disasters in recorded history. The M 7.8 1976 Tangshan earthquake in China killed an estimated 250,000 people, and the 1556 Shanxi earthquake reportedly killed in excess of 800,000. On an annualized basis, earthquakes cause \$5.3 billion in damage in the United States alone [*Federal Emergency Management Agency, 2008*]. Great effort has been put into mitigating the effects of earthquakes in the long and medium terms. We now regularly forecast earthquake probabilities on the scale of a few decades [*Working Group on California Earthquake Probabilities, 2008*] and using these forecasts we generate empirical estimates of ground motions [*Petersen et al., 2008*]. New buildings in the United States are designed to withstand these ground motions and remain life-safe. Meanwhile, public education initiatives such as the Great California ShakeOut seek to increase the proportion of the population that is prepared for an earthquake by raising awareness of disaster preparedness and home retrofit measures.

The elusive aspect of earthquake mitigation has always been the short term. Unlike other natural disasters such as hurricanes or, in many cases, volcanic eruptions, there is no time to evacuate or take shelter once the event has begun. Seismic waves travel at between 10 and 20 times the speed of sound, and the time between the initiation of an earthquake and the onset of shaking at nearby locations is on the order of only a few seconds. Although many different avenues of research have been pursued in the quest for short-term earthquake prediction, there is to date no precursory phenomenon which is both unique to earthquakes and universal to earthquakes. Also, no means exist to translate precursory phenomena into the precise size, location and timing of an imminent event, i.e. an actionable warning. However, recent advances in real-time seismic monitoring have opened up a new avenue for short-term earthquake mitigation: Earthquake Early Warning.

Earthquake Early Warning (EEW) at its most basic relies on detection of seismic activity at one or more stations on a network, and relaying that detection to other sites on the network or to assets in the area protected by the network. Mexico City has been using such a “frontal detection” system for over a decade [*Espinosa Aranda et al., 1995*]. These methods require detecting the damaging S-waves and surface waves from an earthquake before an estimating the intensity of shaking, and thus rely on the fact that the speed of seismic waves, while quite fast, is much slower than the speed of digital communications. The warning afforded by these methods is dependent on the earthquake source being close to the seismometers but far from the assets that require protection. More recent methods detect and characterize the P-waves of earthquakes to estimate the intensity of impending shaking [*Allen, 2004; Wu and Kanamori, 2005a; Wu and*

Kanamori, 2005b; Wu et al., 2006; Allen, 2007], and rely on the fact that the weaker P-waves travel faster than the damaging S-waves and surface waves to enable on-site warning as well as network-based warning. These methods are better suited to regions like California, Japan and Taiwan, where hazardous faults exist in close proximity to population centers.

As a practical matter, EEW systems based on P-waves function by exploiting empirically observed correlations between some property of the early P-wave (amplitude, spectral content, etc.) and the final magnitude of the earthquake. While the statistical significance of the observations, and the observations themselves, continue to be debated [*Rydelek and Horiuchi, 2006; Wolfe, 2006*], these correlations appear to offer some insight into the deeper physics of earthquake rupture. In particular, these observations speak to one of the more hotly-debated topics in modern seismology: whether earthquakes are self-similar, cascading ruptures or whether they scale deterministically from a very early time in the rupture history. The purpose of this dissertation is to address this fundamental question through a combination of observational seismology and kinematic and dynamic modeling.

The first part of this dissertation involves observations of the correlation between the predominant period of early P-wave arrivals and the final magnitude of an earthquake. These observations are part of an ongoing effort to establish an operational EEW system in California using the ElarmS methodology [*Allen, 2007*]. This work establishes empirical scaling relationships for both predominant period and peak amplitude of P-waves with magnitude, and demonstrates the performance of the methodology on two moderate events in the greater San Francisco Bay Area. These observed scaling relationships add to the statistical significance of observations from preceding work, leading to increased confidence in the veracity of these scaling relationships.

In the second part, we seek to corroborate the observations of scaling in P-waves with an independent dataset: kinematic source inversions. Using a large online database of kinematic inversions, we extract the moment release in the early part of the rupture history and correlate it to the final magnitude of each event. We test the correlations against the null hypothesis that earthquakes are purely cascading events and find that this hypothesis can be rejected with a high degree of confidence. This finding, in combination with the observed relationship between the properties of early P-waves and earthquake magnitude, suggests that the early rupture history has some effect on the later evolution of rupture.

The third and final part of this dissertation posits a physical mechanism by which this effect might be mediated. We hypothesize that some property of the early rupture, which we call its “intensity,” imparts more or less energy to the rupture front and respectively either enables or inhibits the rupture to overcome regions of the fault which are unfavorable to rupture.

We simulate this hypothesis using dynamic rupture models in two and three dimensions with imposed heterogeneous shear stress and a rate-and-state friction law. We search over four parameters governing the shear stress distribution, and five governing the friction law using a Genetic Algorithm search, and find plausible values for all but one parameter, suggesting the hypothesized behavior is physically realistic.

Much further work must be done, but these three studies bring us closer to answering the question of whether or not earthquakes are purely cascading phenomena.

2. Toward Earthquake Early Warning in Northern California

Gilead Wurman, Richard M. Allen, and Peter Lombard

Published in the *Journal of Geophysical Research* as:

Wurman, G., R. M. Allen, and P. Lombard (2007), Toward earthquake early warning in northern California, *J. Geophys. Res.*, **112**, B08311, doi:10.1029/2006JB004830.

2.1. Abstract

Earthquake Early Warning systems are an approach to earthquake hazard mitigation which takes advantage of the rapid availability of earthquake information to quantify the hazard associated with an earthquake and issue a prediction of impending ground motion prior to its arrival in populated or otherwise sensitive areas. One such method, Earthquake Alarm Systems (ElarmS) has been under development in southern California and, more recently, in northern California. Event magnitude is estimated using the peak amplitude and the maximum predominant period of the initial P-wave. ElarmS incorporates ground motion prediction equations and algorithms from ShakeMap for prediction of ground motions in advance of the S-wave arrival. The first peak ground motion estimates are available one second after the first P-wave trigger, and are updated each second thereafter for the duration of the event. The ElarmS methodology has been calibrated using 43 events ranging in size from M_L 3.0 to M_w 7.1 which occurred in northern California since 2001. We present the results of this calibration, as well as the first implementation of ElarmS in an automated, non-interactive setting and the results of 8 months of non-interactive operation in northern California. Between February and September of 2006, ElarmS successfully processed 75 events between M_d 2.86 to M_w 5.0. We find that the ElarmS methodology processed these events reliably and accurately in the non-interactive setting. The median warning time afforded by this method is 49 seconds at the major population centers of the Bay Area. For these events the magnitude estimate is within an average of 0.5 units of the network-derived magnitude, and the ground motion prediction from ElarmS is within an average of 0.1 units of the observed Modified Mercalli Intensity.

2.2. Introduction

Earthquake Early Warning (EEW) systems are combinations of instrumentation, methodology and software designed to analyze rapidly an ongoing earthquake and issue real-time information about the hazard to persons and property before the onset of strong ground motions in populated areas. Japan, Mexico, and Turkey currently operate EEW systems, while Taiwan, Italy, Romania and Greece are testing EEW systems [Allen, 2006, and references therein]. Japan's EEW system, which has been providing warnings to a limited group of users, is anticipated to begin widespread public dissemination of warnings in the summer of 2007. The system operating in Mexico is a frontal detection system, which relies on the fact that the largest potential earthquake epicenters are 300 km from Mexico City in the Middle America Trench, such that an array of seismometers between the fault and the city can reliably detect and measure the intensity of an earthquake's S-waves and issue a warning well before those waves arrive at the city [Espinosa Aranda et al., 1995].

In California, the proximity of major faults to population centers limits the utility of frontal detection systems for EEW. Under the conditions found in California a useful EEW system must be able to rapidly and reliably estimate the location, origin time and size of an earthquake based on the P-wave alone. The system must then generate predictions of ground motion at multiple locations of interest and disseminate these predictions in the time between the P-wave arrival and the S-wave arrival. Such systems are being developed in Taiwan [*Wu and Kanamori, 2005a*] and in Japan [*Odaka et al., 2003*] which rely on measurement of the amplitude of the P-wave as a proxy for the magnitude of the earthquake. Such systems are effective for small- and moderate-size events, but are susceptible to saturation in large events. Ground accelerations near the source of large earthquakes saturate at approximately 10-15 m/s², due in part to ground response becoming nonlinear under large stresses.

The Earthquake Alarm Systems (ElarmS) methodology [*Allen and Kanamori, 2003*] has been tested using data from southern California, Taiwan, Japan and the Pacific Northwest of the United States [*Olson and Allen, 2005; Lockman and Allen, 2007*], and uses the maximum predominant period (τ_p^{max}) of the first 1 to 4 seconds of the P-wave as an estimate of earthquake magnitude. The ElarmS methodology has been shown to be effective in these areas for M 3 and larger earthquakes [*Allen and Kanamori, 2003; Lockman and Allen, 2005; Olson and Allen, 2005; Allen, 2006; Lockman and Allen, 2007; Allen, 2007*]. In the process of testing ElarmS in northern California, we find that using both τ_p^{max} and the peak amplitude of the P-wave improves the accuracy of the ElarmS magnitude estimate. We have been testing the effectiveness of the combined methodology since February of 2006 and find that the system estimates the magnitude of earthquakes in northern California rapidly, accurately and reliably.

In addition to incorporating P-wave peak amplitude in the magnitude determination, we have incorporated the attenuation relationships (hereafter referred to as ground motion prediction equations, GMPEs) and algorithms of ShakeMap [*Wald et al., 2005*] into the part of the methodology which predicts ground motions during an event. The GMPEs used by ShakeMap [*Newmark and Hall, 1982; Boore et al., 1997; Wald et al., 1999a; Boatwright et al., 2003*] replace the empirical attenuation relationships developed for southern California [*Allen, 2004; Allen, 2007*]. ShakeMap algorithms [*Wald et al., 1999a*] incorporate individual station corrections to observations as well as scaling of predicted ground motions based on local geology [*Borcherdt, 1994; Wills et al., 2000*] throughout northern California. In addition to making ElarmS ground motion predictions directly comparable to other products like ShakeMap itself, we find the incorporation of these algorithms allows us to generate accurate and timely predictions of ground motion at seismic stations.

2.3. The ElarmS Methodology

Implementing Earthquake Early Warning in northern California presents opportunities not seen in other places for improving the robustness of the ElarmS methodology across different networks. A functional EEW system in northern California must integrate data from both high-gain, broadband velocity instruments and from low-gain, strong-motion accelerometer stations. The system must collect this data over the two networks currently operating in northern California: the Northern California Seismic Network (NCSN) operated by the US Geological Survey, and the Berkeley Digital Seismic Network (BDSN) operated by the UC Berkeley Seismological Laboratory. Within each network, high-gain velocity instruments (channels HHE, HHN and HHZ, which we address as HH henceforth) are more useful for measuring the small ($M < 4.5$) events on which we rely for calibration and routine validation of the method. However, these stations will clip quickly in the event of a nearby major earthquake. Low-gain, strong-motion accelerometers (channels HNE, HNN and HNZ; or HLE, HLN and HLZ, which we address as HN and HL respectively) will remain on-scale longer in the event of a nearby major earthquake, but are of limited use in measuring small events due to low signal-to-noise ratios. Between networks, differences in instrumentation may lead to different behavior within the same channel type (i.e., velocity or accelerometer). All of these differing behaviors must be accounted for by an EEW system which seeks to maximize the amount of usable data in a minimum amount of time.

We use the Earthquake Alarm Systems (ElarmS) methodology developed by *Allen and Kanamori* [2003], with some modifications for the particular problems of northern California. The ElarmS methodology is built of two systems: a single-waveform processing system extracts parameters of interest from a single channel of data, and sends these parameters to an event processing system. The latter integrates output from waveform processing of multiple channels into information about an event's size, time and location, and in fact whether there is an event at all. Given an event's size and location, ground motion predictions are issued for specific locations on a second-by-second basis during the event. Within both of these systems we encounter the need to modify the original ElarmS methodology to account for the specific challenges of northern California data. These will be discussed at length later.

2.3.1. Calibration dataset

Prior to applying ElarmS in a real-time setting, we tested the method on 43 calibration events ranging in size from M_L 3.0 to M_w 7.1 which occurred in northern California since 2001. The calibration events are shown in Figure 2.1. We were restricted from using older events such as the 1989 Loma Prieta and 1992 Petrolia earthquakes in the calibration, because prior to 2001 the NCSN and BDSN networks did not have sufficient station coverage or the appropriate instrument types to measure these events. The

calibration events were used to establish the maximum predominant period vs. magnitude and peak amplitude vs. magnitude relations described below.

2.3.2. Triggering and event location

The first step in the early warning process is to detect an event. This begins with the waveform processing system, which must detect the initial P-wave of an event and issue a trigger at the onset of that P-wave. We use a short-term/long-term average method following *Allen* [1978]. The algorithm is applied to the vertical velocity trace with timescales of 0.5 sec for the short-term average and 5 sec for the long term, and a triggering threshold of 20. Triggering can be accomplished using any real-time algorithm, but cannot be done with any method which requires data after the trigger itself, as such data is by definition unavailable at the time of the trigger.

Consequently, methods such as auto-regressive pickers [*Sleeman and van Eck*, 1999] and pickers based on wavelet transforms [*Zhang et al.*, 2003], while more precise than a simple short-term/long-term average method are not practical for this application. This also means generally that any filter applied to the data must be causal.

When the first station triggers, the event processing system will provisionally locate the event beneath that station. When a second station triggers the provisional location moves to a point directly between the two stations, based on the timing of the arrivals. Once trigger times are produced at three or more locations, the event location and origin time is estimated using trilateration and a grid-search algorithm to find the optimal solution. Although a depth can be estimated using more stations or more sophisticated algorithms, this is unnecessary for the geologic setting of northern California, where most events nucleate at less than 20 km depth [*Hill et al.*, 1990]. We currently fix the depth of the event to be 8 km.

Based on the estimated event location and time, warning times can be calculated for any geographical locations of interest. These warning times are based on a move-out speed of 3.75 km/s, which is determined from observations of the onset times of significant ground motions in southern California. Again, although more sophisticated methods exist for the estimation of time until significant shaking, when one considers the computational requirements for greater sophistication against the need for rapid processing and notification, this simple move-out speed seems sufficient for the purpose of estimating the warning time.

2.3.3. Magnitude from predominant period

The ElarmS methodology rests largely on the use of the maximum predominant period (τ_p^{max}) within the first 4 seconds of the P-wave as an indicator of the size of the event [*Allen and Kanamori*, 2003; *Olson and Allen*, 2005]. The predominant period, τ_p of a single vertical channel (HHZ, HLZ or HNZ) is calculated in real time using the iterative relation

$$\tau_{p,i} = 2\pi \sqrt{\frac{X_i}{D_i}} \quad (2.1)$$

where $X_i = \alpha X_{i-1} + x_i^2$ and $D_i = \alpha D_{i-1} + (dx/dt)_i^2$. The constant α is a smoothing constant, and x_i is the ground velocity of the last sample. Because both velocity sensors and accelerometers are used, the accelerometer traces must be integrated to velocity before τ_p can be calculated. In addition, a causal 3 Hz low-pass Butterworth filter is applied iteratively to the velocity data [Allen and Kanamori, 2003]. This calculation is done by the waveform processing system, and the maximum value of τ_p within the first 4 seconds of the P-wave arrival is recorded and sent to the event processing system, which uses it to estimate magnitude according to a predetermined relationship.

Our initial attempts to establish a relationship between magnitude and τ_p^{max} were frustrated by noise in the low-magnitude data ($M < 4.5$). This problem led us to adopt two significant additions to the ElarmS methodology. The first of these is a criterion for the disqualification of S-wave data. Part of the low magnitude scatter was due to many small events being located close to our stations in the San Francisco Bay Area. As a result, the S-wave arrival occurs within 4 seconds of the P-wave arrival, and since S-waves generally have longer periods than the associated P-waves, it is the S-wave τ_p which gets recorded as τ_p^{max} . A simple criterion based on an S-minus-P move-out of 1 second per 8 km eliminates these false signals and cleans up the data somewhat, though we do apply a minimum S-minus-P time of 1 second, based on the assumption that the event is 8 km deep.

In addition to this S-wave criterion, we chose to incorporate a second criterion for the exclusion of data, based on the signal-to-noise ratio (SNR) of each waveform. As this was a particular problem for low-gain accelerometers (HL and HN channels), we chose to treat each channel type separately. The absolute noise level is calculated as a very long-term average from inter-event data, and is frozen when a trigger is detected. From the time of the trigger until the event is over, the signal level is calculated using a 0.05-second short-term average [Allen, 1978], and the ratio of these two is the SNR. In principle the higher we require the SNR to be, the better our results. However, we must consider the need for fast measurements as well as good ones, and the greater SNR we require, the fewer measurements of the first second of the P-wave will be admitted. By weighing the reduction in scatter of small magnitude τ_p^{max} against the number of excluded τ_p measurements in the first second of data, we arrive at the optimal minimum SNR: 100 for HH channels and 200 for HL and HN channels.

The results of calibrating τ_p^{max} vs. magnitude are shown in Figure 2.2. Note that low-gain accelerometer data still shows a significant scatter in

spite of the two added criteria. We are investigating the root cause of this scatter, but presently the HN channels (Figure 2.2c) exhibit the largest scatter, and we have provisionally removed them from the τ_p^{max} determination until this can be resolved. The best fit relationship between τ_p^{max} and magnitude is

$$M = 5.22 + 6.66 \cdot \log_{10}(\tau_p^{max}) \quad (2.2)$$

using only the HH and HL channels. This relationship is plotted in Figure 2.2. In determining magnitude from τ_p^{max} for any new event we only use data from HH and HL channels to be consistent with the calibration of this relationship.

2.3.4. Magnitude from P-wave peak amplitude

While we have succeeded in reducing the scatter in measurement of τ_p^{max} at low magnitudes, the scatter is still sufficient to present us a problem in discriminating between small non-hazardous events and large hazardous events. Because of this scatter, there is a potential to misidentify a small event as a large one, leading to a false alarm. This is of critical importance in many early warning applications, as a high incidence of false alarms will drastically reduce the credibility and utility of the warnings. This is especially true in applications where the cost of false alarms is high, such as industrial process interruption. In order to further improve this discrimination, we have added a second, independent estimate for rapid magnitude determination. Using a method similar to that of *Wu et al. [2006]*, we calculate the peak amplitude of the P-wave, scaled by the logarithm of the epicentral distance. As with τ_p^{max} , we glean the peak amplitude from the first 4 seconds of the vertical record. *Wu et al.* used the peak displacement, P_d of the P-wave, but we chose to analyze displacement, velocity and acceleration for each channel type independently. In theory, the displacement record has longer periods than the acceleration or velocity records, and will be less susceptible to random high frequency excursions. For velocity instruments (HH channels) this holds true, and measuring the peak amplitude in displacement yields the lowest error. However, for accelerometer channels (HN and HL), the act of numerically integrating twice (from acceleration to velocity and then again to displacement) introduces errors to the point where using the velocity record rather than displacement yields a better magnitude estimate.

We also investigated the merit of using between 1 and 5 seconds of P-wave data for determining P_d or P_v (peak velocity, for HN and HL channels). Using less than 4 seconds yielded greater errors, and between using 4 and 5 seconds there was little difference in performance (4 seconds performed slightly better for HH, slightly worse for HL and HN channels). We chose to use 4 seconds for the sake of internal consistency with our τ_p^{max} measurements, which also use 4 seconds of P-wave data.

The results of calibrating P_d and P_v (which we henceforth abbreviate $P_{d/v}$) vs. magnitude are shown in Figure 2.3. The amplitudes are plotted as a function of magnitude, after being scaled to an epicentral distance of 10 km using the best-fit relations in Eqs. 2.3 through 2.5 below. These plots do not show nearly the scatter at low magnitudes that the τ_p^{max} vs. magnitude plot does in Figure 2.2. However, the $P_{d/v}$ of the largest (M_w 7.1) event is significantly lower than predicted. This is due to the fact that this event incorporates data from more distant stations than is normally allowed, as will be discussed in detail later. Due to this effect, the $P_{d/v}$ measurements for the M_w 7.1 event were not used in the best fit lines plotted in Figure 2.3.

Although the variability of the HN data seen in τ_p^{max} is visible to a lesser degree in P_v , (Figure 2.3c) the data are not unusable. However, we chose to fit HL and HN data separately to minimize the error of measurements on the HL channels. The best fit relationships between magnitude and $P_{d/v}$ are

$$M = 1.04 \cdot \log_{10}(P_d) + 5.16 \cdot \log_{10}(R) + 1.27 \quad (2.3)$$

(HH channels)

$$M = 1.37 \cdot \log_{10}(P_v) + 4.25 \cdot \log_{10}(R) + 1.57 \quad (2.4)$$

(HL channels)

$$M = 1.63 \cdot \log_{10}(P_v) + 4.40 \cdot \log_{10}(R) + 1.65 \quad (2.5)$$

(HN channels)

where R is the distance in kilometers from the station to the epicenter. When the source of scatter in the HN data is found and controlled for, it may be beneficial to unify the relationships for HL and HN channels. Since the HH channels use P_d rather than P_v the relationship for HH channels must remain separate from the other two.

2.3.5. Data integration and magnitude determination

The waveform processing system sends any new information available to the event processing system every tenth of a second for four seconds after a trigger. This includes the maximum value of τ_p^{max} or $P_{d/v}$ only if that maximum has changed since the last tenth of a second. The τ_p^{max} data is accompanied by the value of the SNR at the time of the measurement. This low data volume has the advantage of being easily transmissible over existing station telemetry, so that the waveform processing system can potentially be implemented at each station independent of the rest of the network. The advantage of this approach, in turn, is that the waveform processing happens much sooner and much more reliably, as there is no delay for telemetry of data over the network, and no risk of data dropout leading to errors in processing. Instead, the large volume of data being

produced by the sensors is reduced on site to a few parameters of interest which can be cheaply transmitted over the network.

The event processing system gathers the transmitted data from the waveform processing systems at each station within 100 km of the estimated epicenter. This is the distance within which frequency-dependent attenuation (Q) has a minimal effect on the predominant period measurement [Allen and Kanamori, 2003]. For the two largest calibration events, the M_w 6.5 San Simeon earthquake and the M_w 7.1 earthquake in the Gorda Plate, this cutoff distance is increased to 150 km and 200 km respectively, due to the lack of stations within 100 km of these events. We justify this in particular for the Gorda Plate event by asserting that the intervening crust between the event and the stations is mostly oceanic, and has higher Q than continental crust [Vera et al., 1990]. The system integrates the data from the stations once per second to determine a magnitude estimate for the event as it progresses. Each time a new maximum τ_p^{max} or $P_{d/v}$ value is reported, the event processing system checks it for validity by examining whether the S-wave may have arrived at that station, as described earlier in this section. It also checks that the SNR at the time of a τ_p^{max} measurement exceeds the minimum required value. If any of these checks fail, the event processing system ignores that measurement and proceeds as if it was never reported.

The event processing system makes one more check, in which it looks for an indication that a given channel has clipped. This indication is actually given by the waveform processing system in the form of a negative SNR beginning when the channel's output first exceeds a particular threshold, and extending for a fixed duration after the last sample which exceeds that threshold. This duration represents the time required for the channel to recover from the clipping event and become usable. The clipping threshold and recovery time vary from channel to channel, and are encoded in the waveform processing system at each station, so the event processing system does not know anything about the value of the data, only whether it has clipped. If the event processing system receives a clipping indication, it immediately stops updating information from that station for the duration of the event. The τ_p^{max} value at the time of clipping is recorded as the final τ_p^{max} for that station, and the $P_{d/v}$ value for the station is stricken. The reason for treating the two estimates differently is that we often find that the time at which τ_p^{max} is taken is not the same time as the peak amplitude of the P-wave, so the τ_p^{max} value before the clipping occurred is still potentially valid. In contrast, the fact that the sensor has clipped means a priori that the previous $P_{d/v}$ value has been exceeded, and is therefore invalid. In this respect τ_p^{max} is more robust, as it can tolerate clipping of a channel and still represent a valid estimate.

If the data passes all the checks, the quality of the data can be reasonably assured, and the event processing system uses the updated

information to produce a magnitude estimate for the event. It takes the $\log_{10}$ average of τ_p^{max} from each available channel, and calculates a magnitude from the average value. The results of magnitude estimation for the calibration events using τ_p^{max} alone are shown as gray triangles in Figure 2.4. Note the significant scatter in the magnitude estimate below $M \approx 4.5$, consistent with the calibration of τ_p^{max} vs. magnitude from Figure 2.2.

The event processing system also takes the average value of $P_{d/v}$ from each station and calculates a magnitude from that average value. The results for the calibration events are shown as gray squares in Figure 2.4. Note the comparatively low magnitude assigned to the largest event in the calibration dataset, consistent with Figure 2.3. As discussed earlier, this is the result of incorporating data from more distant stations for this event. However, it highlights a potential limitation of the $P_{d/v}$ estimate. The $P_{d/v}$ estimate is susceptible to saturation near the fault for very large events. This is because at fault-normal distances less than the length of the rupture the distance to the farthest point of the rupture is significantly greater than the distance to the nearest point. As a result, the effective distance between the station and the rupture (i.e., the average distance between the station and all points on the rupture) is greater than the actual fault-normal distance, leading to lower P-wave amplitude than predicted for a given epicentral distance.

The value of τ_p^{max} does not appear to be susceptible to this effect [Olson and Allen, 2005], but is much more susceptible to noise pollution at lower magnitudes than peak amplitude measurements. Thus the two estimates of τ_p^{max} and $P_{d/v}$ are particularly complementary, with the strengths of one compensating for the weaknesses of the other, and using some combination of the two magnitude estimates from τ_p^{max} and $P_{d/v}$ produces a more robust single estimate. Currently, the two estimates are combined in a linear average, the results of which are shown as black circles in Figure 2.4. Note the superior performance at both ends of the magnitude scale as a result of this combined approach. The large events are not underestimated, and the scatter in the small events has been greatly reduced.

A more sophisticated scheme may be conceivable for the combination of the τ_p^{max} magnitude with the $P_{d/v}$ magnitude. In particular, since we are interested in the low-magnitude performance of $P_{d/v}$ and the high-magnitude performance of τ_p^{max} , it makes sense to consider a progressive weighting scheme in which the latter is more heavily weighted at low magnitudes and the former more at high magnitudes. We investigated a scheme by which the weighting changes linearly with the magnitude of the event, but found that the data does not bear out the use of such a scheme. At this time the simple linear average appears to be as good as any weighted average, so we use the linear average.

2.3.6. Ground motion prediction

The final step in an early warning system is to predict the severity of imminent ground motions from an ongoing earthquake and to issue warnings based on those predictions. We do not address the question of when and how to issue warnings. For a discussion of this aspect of EEW, see the work of *Brown et al.* [2009]. For this study, we observe that the $1-\sigma$ error in magnitude estimate reduces to a reasonable level (0.5 magnitude units) when 4 seconds of data are available from 4 channels. We define this criterion of 4 seconds of data in 4 channels as the “alarm time” for the purposes of performance evaluation in the next section. However, we arrive at this definition somewhat arbitrarily, and different users would require different levels of uncertainty or timeliness, depending on their tolerance for false or missed alarms [*Brown et al.*, 2009].

ElarmS is capable of generating ground motion predictions through the incorporation of algorithms from ShakeMap [*Wald et al.*, 2005]. These algorithms, which have been developed for and tested extensively in California, incorporate empirically-derived GMPEs [*Newmark and Hall*, 1982; *Boore et al.*, 1997; *Wald et al.*, 1999a; *Boatwright et al.*, 2003], as well as geological amplification correction and corrections for site conditions at seismic stations [*Borcherdt*, 1994; *Wills et al.*, 2000]. The ground motion predictions are initially calculated using only the estimated magnitude and location of the event, as no observations of peak ground motion are yet available. We use the GMPE for the given magnitude to compute the predicted ground motion on a regular grid of points with a spacing of 0.1° around the source. The prediction at each point is then corrected for local geological effects. The result is a coarsely spaced grid of points with predictions of peak ground motion based solely on the magnitude and location of the event. This grid can be interpolated to create predictions at finer resolution, and to generate predictions for discrete locations of interest, such as urban centers or seismic stations.

As the event progresses and the S-wave field expands outward from the source, peak ground motion observations become available at each station. The observations are first corrected for the site condition at the station, and then the GMPE curve, based on magnitude and location, is linearly scaled up or down to best fit the corrected observations. The resulting equation is used to generate ground motion predictions on a regular grid as before, with the addition of grid points representing the individual station observations available at the time. This irregular grid is interpolated to produce a finer, regular grid of peak ground motion incorporating magnitude, location and station observations. This grid is predictive at all points ahead of the S-wave front. The process is similar to that used to produce ShakeMaps after an earthquake, but here it is done once per second. Initially there is very little information to incorporate and the ground motion predictions are correspondingly rough, but with each second that passes the information

becomes more complete and the ground motion predictions are refined in real time. ElarmS produces predictions of PGA and PGV at all points, which are combined to produce a prediction of Modified Mercalli Intensity (MMI) using the relationship of *Wald et al.* [1999b].

2.3.7. Simulating ElarmS

It is not practical to implement the ElarmS methodology online as outlined in the first part of this section without first testing it offline to ensure its functionality. This is because a full implementation requires the investment of time and money to emplace the waveform processing system at each station in the network. Therefore, we test the performance of the methodology offline using a program that simulates the causality of information after the event has completed. While this simulation may differ somewhat from the final implementation of ElarmS, the behavior of the methodology will not change appreciably from the results of the simulation. Henceforth, when referring to "ElarmS" we refer to the simulation unless otherwise stated.

2.4. ElarmS Performance

Since February 2006, we have been operating ElarmS automatically following every event of M 3.0 or larger in northern California. This processing is initiated 10 minutes after notification of a new event, in order to allow the requisite data to be collected at the network data center for retrieval. The processing is performed automatically with no human input or oversight. We have been using the results of this automatic processing to make improvements to the ElarmS methodology, and consequently it is necessary on occasion to re-process these events after the fact, when a significant change is made in the methodology. This reprocessing is prompted by a human operator, but without any added input from the operator. The process is identical to the automatic processing and uses the same data which was gathered 10 minutes after each event. We call this "non-interactive" processing, and we use it to indicate how a real-time implementation of ElarmS might perform.

2.4.1. Performance of non-interactive processing

Between February and September of 2006, a total of 85 instances of non-interactive processing were initiated. Of these, one is a duplicate event, a result of the email notification system posting an update to an existing event. One instance was a false event. This was not the result of a false detection by ElarmS, but of an erroneous email notification.

The geographic distribution of the remaining 83 events is shown in Figure 2.5. Of these, one event was offshore Mendocino, with no stations within 100 km of the source. This is the cutoff distance for usable stations in the ElarmS methodology, so the event produced no output. Seven events suffered errors as a result of maintenance of the operating system. Five of

these occurred consecutively, due to the extraordinary misfortune of an update to the operating system coinciding with a temporal (not spatial) cluster of small events (all $M_L \leq 3.6$). The system maintenance prevented the acquisition of data 10 minutes after the earthquake. Acquiring data at a later time invalidates the “non-interactive” procedure, so these events are not considered in this analysis.

The remaining 75 events range in magnitude from M_d 2.86 to M_w 5.0. The results of magnitude estimation for these 75 events are presented in Figure 2.6. This figure shows the magnitude errors (with respect to network-based magnitudes, usually M_w or M_L) produced by ElarmS at three different times for each event. The initial magnitude error (Figure 2.6a) refers to the magnitude estimation based on only the first second of P-wave data at the first station or stations to detect the event. This is the earliest possible magnitude determination, which can be used to give the maximum warning time. The initial magnitude has a significant scatter ($\sigma = 0.72$ magnitude units) due to its reliance often on a single station's data.

Figure 2.6b shows the errors at “alarm time”, which we define as in the previous section to be the time at which at least four seconds of P-wave data are available from at least four different channels. The magnitude error at this time is considerably less than in the first second ($\sigma = 0.54$ magnitude units). There are fewer events represented in this plot (66 events vs. 75 in Figure 2.6a), because not all of the events are ever detected in enough channels to meet the alarm criteria. This is primarily due to the weak signal from small ($M \approx 3$) events, and in some cases results from a lack of enough stations within 100 km of the epicenter.

Figure 2.6c shows the error in the final magnitude determination for events that met the alarm criteria, using all available data from stations within 100 km of the source. The scatter has decreased slightly ($\sigma \approx 0.48$ magnitude units) due to the incorporation of more station information. In all three of these plots, the magnitude estimate is biased slightly downward (mean of -0.57 magnitude units in the initial estimate, -0.13 at alarm time and -0.02 in the final magnitude estimate). This is due to events beyond the physical edge of the network, which can be mislocated by tens of kilometers due to poor azimuthal coverage. This does not affect τ_p^{max} -based magnitude estimates, but $P_{d/v}$ -based magnitude estimates are strongly affected by epicentral distance errors. Location errors can also cause the system to set the S-wave arrival time earlier than the true arrival time because it considers some events to be closer to the station than it is. This causes the system to discard valid data when measuring both τ_p^{max} and $P_{d/v}$, because it considers the signal contaminated by the S-wave. This biases the estimates downward because it prevents both τ_p^{max} and $P_{d/v}$ values from being revised at later times, and these revisions are always upward.

Figure 2.7 summarizes the error in ground motion prediction at seismic stations in the NCSN and BDSN networks for all events. The logarithm of the

observed ground motions is subtracted from the logarithm of the predicted ground motions for PGA (Figure 2.7a,b) and PGV (Figure 2.7c,d). An error of 1 signifies over-prediction by a factor of 10. Figure 2.7a and c show the errors in the first second of data, and Figure 2.7b and d show the errors at alarm time. At alarm time the $1-\sigma$ error in PGA and PGV is approximately a factor of 4 (0.6 log units). For the MMI errors (Figure 2.7e,f) no logarithm is necessary as MMI incorporates logarithms of PGA and PGV [Wald et al., 1999b]. The number of predictions in the center bin (less than ± 0.17 MMI unit error) is off scale in Figure 2.7e and f. There are 435 observations in the center bin in the first second (Figure 2.7e) and 398 at alarm time (Figure 2.7f). The scatter in ground motion prediction is significantly reduced by waiting for the alarm condition to be met ($\sigma=0.42$ MMI units in the first second, versus $\sigma=0.08$ at alarm time). Figure 2.7e has a positive bias (0.12 MMI units). This is because the ShakeMap MMI scale ranges from 1 to 10, and for many of the smaller events MMI cannot be significantly under-predicted simply because the actual MMI is only 1 or 2. There is also a slight negative bias (-0.01 MMI units) in Figure 2.7f reflecting the same effects as seen in Figure 2.6. The errors reported in Figure 2.6 and Figure 2.7 are valid only for events in the validation dataset. We cannot evaluate the error for larger earthquakes, as none occurred in the time of the study.

The times between event origin and detection and the achievement of the alarm condition for all the events are summarized in Figure 2.8. Initial detection occurs an average of 8.0 ± 4.8 seconds after event origin (Figure 2.8a). The alarm condition is reached an average of 14.9 ± 4.6 seconds after origin (Figure 2.8b). These results are presented in Figure 2.9 in terms of time until largest ground shaking at three major Bay Area metropolitan centers: San Francisco, San Jose and Oakland. The median warning time in these cities at initial detection is 56 seconds in San Francisco (Figure 2.9a) or Oakland (Figure 2.9e), and 48 seconds in San Jose (Figure 2.9c). If we wait for the alarm condition to be reached, the median warning times reduce to 39.5 seconds in San Francisco (Figure 2.9b), 40 seconds in San Jose (Figure 2.9d) and 39 seconds in Oakland (Figure 2.9f). This analysis does not show the warning time for any future earthquakes, but the distribution of event locations in Figure 2.5 does coarsely reflect the potential locations of future large earthquakes. For a more detailed analysis of warning time for potential damaging earthquake scenarios, see *Allen* [2006].

The two event parameters which have not been discussed are the epicenter estimates and the origin time estimates. When only one station has triggered, ElarmS assumes the event origin time and epicenter correspond to the time and location of that first trigger. When two stations have triggered, the epicenter is located at a point between the first two triggered stations, based on timing. Consequently location and origin time errors are significant when fewer than three stations are used. However when three or more stations are used we find the error is small, and by the

time the alarm condition is met both of these estimates have insignificant error. This characterization does not hold as well for events beyond the edge of the network. At alarm time, the mean absolute error in epicenter location is 13.7 ± 23.4 km and the mean absolute error in origin time is 2.3 ± 3.1 seconds, including events beyond the edge of the network. For both of these measures, the mean error is within a standard deviation of zero.

2.4.2. Two Bay Area scenario events

Among the 75 events processed non-interactively by ElarmS, two moderate events represent likely hazardous earthquake scenarios for the Bay Area (Figure 2.10), and thus provide some insight into what can be expected of ElarmS. For these two events we use M_L as a reference, even though M_w values exist for both. This is because M_L is sensitive to the same frequencies (~ 1 -2 Hz) as ElarmS, and because M_L is more directly related to the severity of the event in terms of damage to persons and property.

The first event is a M_L 4.7 event near Gilroy, CA on 15 June, 2006. This event is located near the southern Calaveras fault, in a geographic location where a Calaveras or Southern Hayward fault rupture might nucleate (Figure 2.10). Figure 2.11a shows the magnitude estimate for this event as a function of time in relation to the arrival time of significant shaking at San Francisco, Oakland and San Jose (vertical lines). The time at which the alarm condition was reached is also plotted, and the dashed horizontal line represents $M = 4.7$, the actual local magnitude of the event.

Figure 2.11b shows the error in predicted PGA vs. time, over all stations which have not yet reported peak ground motion at the given time. The solid line is the mean error at each time, and the dashed lines are ± 1 - σ error margins. Figure 2.11c shows the error in predicted PGV vs. time for the event in a similar fashion. Both of these factor into the predicted MMI for the event [Wald et al., 1999b].

Figure 2.11d shows the error in predicted MMI vs. time in a manner similar to Figure 2.11b and c. After 13 seconds the only stations which have not reported peak ground motion are those far from the fault, which experience an intensity of 1. The predicted MMI at all these stations is 1 (the ShakeMap algorithm does not produce MMI less than 1), and this is reflected in Figure 2.11d by the fact that the mean and ± 1 - σ lines converge to zero after 13 seconds.

Figure 2.12 shows the predicted peak ground shaking from ElarmS (which we call the "AlertMap") for 7 seconds following event detection for comparison with the ShakeMap for this event (Figure 2.12h). The time since event origin and magnitude estimate are given above each AlertMap. Stations are plotted as small white symbols in the same fashion as the maps in Figure 2.1 and Figure 2.5: triangles for velocity sensors, inverted triangles for accelerometers and diamonds for a collocated installation. Stations which have triggered are plotted as larger gray symbols, and stations which are experiencing peak ground motions are plotted in black. When the peak

observations become available, the stations are colored according to their observed peak MMI, following the scale at the bottom of the plot. The circular contours radiating from the epicenter represent the estimated time until the onset of the largest ground motions at all locations, based on the current epicenter location and a move-out speed of 3.75 km/s. The color field in Figure 2.12b-g represent estimated peak MMI following the scale at the bottom of the plot, which is the same as for the ShakeMap in Figure 2.12h. At all points outside the 0-second warning time contour, this peak MMI estimate is predictive.

The first AlertMap in Figure 2.12a, 3 seconds after event origin, represents the initial detection time for this earthquake. There is no magnitude estimate yet, since ElarmS requires a full second of P-wave data before making the initial estimate. The larger gray station is the first station to trigger for this event, and the red star represents the epicenter, currently located at the station as described above. The initial magnitude estimate of M 5.0 is available one second later, 4 seconds after the origin (Figure 2.12b). With the magnitude estimate ElarmS begins predicting ground motions based only on the GMPE. The epicenter has also been relocated at this time due to a second station triggering. As described above, the epicenter is now located directly between the two triggered stations based on the trigger times. At 5 seconds after origin (Figure 2.12c), a third station triggers and from this point forward the location is fit to the trigger times using a grid search. At this time, data is being collected from 4 channels (two channels at the station southeast of the epicenter, one at each of the other two stations), so 4 seconds later (i.e., 9 seconds after origin) the alarm condition will be reached. At 6 seconds (Figure 2.12d) two more stations have triggered, but by now the epicenter location is good and does not move noticeably.

At 7 seconds after origin (Figure 2.12e), the magnitude estimate has dropped to 4.7, and the first peak ground motion observation is available. At this time, the GMPE curve is biased to pass through that single observation, which is why the predicted MMI field for the event changes so drastically between Figure 2.12d and e. The station southeast of the epicenter is plotted in black to indicate that it is currently experiencing peak ground motion, consistent with its location within the 0-second warning time contour. At 8 seconds (Figure 2.12f) two more stations have triggered and another has entered the peak ground motion window. At 9 seconds (Figure 2.12g) the alarm condition is reached since the fourth channel triggered at 5 seconds, and the magnitude estimate is 4.3, only 0.4 magnitude units below the actual M_L . At this time, 3 additional stations have triggered, bringing the total to 13 triggered channels at 10 different stations. Another station has entered the peak ground motion window, and the station southeast of the epicenter has reported its peak ground motion observation. When there are multiple observations, ElarmS biases the GMPE curve to best fit the available

observations. In this case, this results in a slight increase in the predicted ground motion over Figure 2.12f.

The reason for the low intensities far from the event in the final AlertMap (Figure 2.12g) versus the ShakeMap (Figure 2.12h), is that the ShakeMap incorporates peak ground motion observations from stations on the Peninsula and in the East Bay. At 9 seconds after the origin, the S-wave front has not arrived at many of these stations, so that information is not used to bias the GMPE curve in the AlertMap. However, in Figure 2.11d it is apparent that the ElarmS ground motion predictions at 9 seconds are accurate to within a standard deviation of 0.3 MMI units of the actual observed intensities. The vertical lines in Figure 2.11 represent the arrival of the largest ground motions at the three major urban centers in the Bay Area. San Jose experienced peak ground shaking only 12 seconds after event origin, meaning San Jose would have had about 3 seconds warning time in this event, not considering telemetry and dissemination delays. However, Oakland and San Francisco would have had 20 and 22 seconds of warning, respectively for this event. These warning times depend primarily on the disposition of stations around the epicenter, so they would be comparable for a magnitude 7 event. In this case, the distance between the epicenter and the major cities is comparable to that of the M_w 6.9 Loma Prieta Earthquake. Thus, in the case of a Loma Prieta repeat, we would expect comparable warning times in the major cities.

The second scenario event is a M_L 4.7 event near Santa Rosa on 2 August (local time), 2006. This event is located near the Rodgers Creek fault, near possible epicenter locations for a southward-rupturing Rodgers Creek/Hayward fault event (Figure 2.10). Figure 2.13 and Figure 2.14 show the history of this event in the same manner as for the Gilroy event.

Initial detection of this event occurs 3 seconds after event origin, as shown in the AlertMap in Figure 2.14a. The epicenter at this time is collocated with the only triggered station. At 4 seconds (Figure 2.14b), two other stations have triggered, so the epicenter is located using a grid search method. In addition, one of the stations has both an accelerometer channel and a velocity channel, bringing the triggered channel count to 4. Thus the alarm condition will be reached in 4 seconds, at 8 seconds after origin. The initial magnitude estimate at this time is 5.8, over one magnitude higher than the actual size of the event. The high magnitude estimate in turn causes the ground motion predictions to be high, as these are produced using the GMPE alone, with no station observations to bias the curve. This is reflected both in the AlertMap, which shows significantly higher MMI than the ShakeMap (Figure 2.14h), and in Figure 2.14b, which shows that the first MMI predictions exceed actual observations by as much as 2 MMI units.

The comparatively large magnitude error highlights the utility of waiting for more data to become available rather than issuing the alarm immediately, based on information from a single station only. In this case,

the large error is due to the first triggered station being a strong-motion accelerometer. Most of the stations to the north of the Bay Area are strong-motion accelerometers, which are susceptible to noise pollution below $M \approx 5$. For large earthquakes this is not a problem, but in smaller events high-gain broadband velocity sensors yield superior data.

At 5 seconds (Figure 2.14c) after origin the magnitude drops to 4.3 due to the incorporation of one second of data from the three channels which triggered in the previous second. The magnitude estimate is now 0.4 units lower than the actual M_L for this event, but the error is nearly a third of that in the previous second, again suggesting that it is better to wait one second for multiple stations to provide data rather than relying on a single-station estimate. At 6 seconds after origin (Figure 2.14d) the first peak ground motion observations become available, and the GMPE curve is biased to minimize the errors at these stations. Two more stations to the southwest have triggered at this time, leading to a small revision in the epicenter location. At 7 seconds after the origin (Figure 2.14e) the magnitude estimate drops slightly to 4.2. Two more stations have triggered, but the location does not change noticeably after 6 seconds from origin. The station directly south of the epicenter now reports an additional peak ground motion observation, further informing the ground motion predictions at this time.

At 8 seconds after origin (Figure 2.14f) the alarm condition is reached. Two more stations have triggered in this second, bringing the total to 10 channels at 9 stations. The magnitude estimate is 4.2, half a magnitude lower than the actual M_L of 4.7. However the peak ground motion predictions are biased up by the available station observations of peak ground motion, and match both the ShakeMap (Figure 2.14h) and the final observed peak ground motions at the stations outside the 0-second warning contour. Figure 2.14b shows that the MMI predictions are accurate to within a standard deviation of 0.3 MMI units at alarm time. For consistency with Figure 2.12 the AlertMap at 9 seconds, one second after the alarm time, is shown in Figure 2.14g. There is no noticeable change from Figure 2.14f, other than one additional station having triggered.

When comparing the ElarmS AlertMap for this event (Figure 2.14f) with the ShakeMap (Figure 2.14h) the performance appears better at alarm time than for the Gilroy event, even in light of the low magnitude estimate. The two maps are almost identical in terms of peak intensities, though the AlertMap does under-predict the intensity near the epicenter as a result of the low magnitude estimate. At alarm time, both San Francisco and Oakland have 11 seconds until the arrival of the largest ground shaking. The magnitude estimate is low, and will only rise to 4.6 at 13 seconds after the origin (Figure 2.11b) leaving 6 seconds of warning for San Francisco and Oakland, but this is not an issue when considering ground motion predictions, which are accurate at this time. San Jose experiences its largest ground motions 37 seconds after origin, so even with the additional 5 second

delay for the magnitude estimate to rise, it still has 24 seconds of warning in this instance.

2.5. Improving ElarmS Performance

The performance of ElarmS in the non-interactive processing arena is promising. At alarm time the 1- σ magnitude error for these events is 0.5 magnitude units, which is consistent with the results using the calibration events, and also consistent with previous work with the ElarmS methodology in southern California [Allen, 2007]. This indicates that we now have a good understanding of how ElarmS behaves in a real setting, at least for earthquakes less than $M \approx 5$, and suggests that we can expect the same behavior in the future in northern California and in other locations as well.

Although the results are largely favorable, there are a few things to consider for full online implementation. For one thing, the warning times obtained in the non-interactive processing are maximum warning times, as they do not take into account telemetry, processing or dissemination delays. The processing time delay should not exceed one second if we are to run ElarmS with continuous updates every second. In terms of station telemetry, the actual transmission delays are negligible [Uhrhammer, personal communication, 2006], but currently data is telemetered in up to 10-second-long packets for stations in the BDSN and NCSN networks [Neuhauser, personal communication, 2006]. This packetization can be reduced to 1 second or less, but the overhead associated with transmitting each packet gets proportionally larger as the packet itself gets smaller. Finally, these warning times do not account for delays in disseminating the data and taking action at the user end. We cannot quantify these delays as no dissemination system exists at this time. However, the results reported here show that a dissemination delay of less than 1 second is ideal. When such a system is designed, the minimization of dissemination delays must be a primary design goal.

One way to significantly improve warning times is to improve the disposition of seismometers around northern California. At present the broadband velocity and strong-motion accelerometer stations in the NCSN and BDSN networks are distributed somewhat unevenly, and not always in optimal locations for observing an earthquake near its epicenter. In particular, most of the stations are located in and around the Bay Area, with comparatively few stations along the northern coast of California. The mitigating circumstance here is that most of the population of northern California lives in the Bay Area, and the preponderance of faults capable of major earthquakes is in and around the Bay Area, where most of the instruments are located. However, the comparative lack of instruments along the northern portions of the San Andreas Fault mean first that any people living north of the Bay Area would receive limited benefits from

ElarmS or a similar EEW system. The second consequence of this station distribution is that an event nucleating on the northern portions of the San Andreas Fault and propagating southward would take a long time to achieve the alarm condition due to the dearth of stations in the vicinity of the epicenter. This means drastically reduced warning times for the Bay Area and likely a larger error margin due to the comparatively few measurements that would be incorporated in the event estimates. Given that a southward-propagating rupture on the northern San Andreas Fault may cause as much as \$90 - \$120 billion in losses in the Bay Area [Kircher et al., 2006] it is worthwhile to instrument the northern reaches of the San Andreas more thoroughly.

Even within the Bay Area there is some room for improvement. There are few broadband or strong-motion stations around the southern segment of the Hayward Fault, which is considered the most hazardous fault in the region [Working Group on California Earthquake Probabilities, 2008]. Given this fault's proximity to the major urban centers in the Bay Area, as little as one or two seconds' additional warning could make a significant difference. A few well-placed seismometers along the Hayward Fault would go a long way toward attaining those extra one or two seconds.

2.6. Conclusions

The ElarmS methodology incorporates two independent measurements to estimate the magnitude of an event: the peak amplitude of the P-wave and its maximum predominant period, both within 4 seconds of the onset of the P-wave. These measurements are controlled both for sufficient signal-to-noise ratio and to prevent pollution by the S-wave arrival. The magnitude estimate produced by ElarmS is a linear average of the estimates determined from peak amplitude and maximum predominant period. The epicenter of an event is located using a grid search algorithm based on the trigger times at three or more stations. The epicenter and magnitude estimates are used to generate predicted ground motions at all points around the epicenter, using algorithms similar to ShakeMap. The predictions are updated each second based on updated epicenter and magnitude information. As observations of peak ground motion become available at stations near the epicenter, the prediction is corrected to conform to these observations, further refining the prediction.

The ElarmS methodology has been applied in an offline simulation to every event greater than M 3 in northern California since February of 2006. The methodology has been applied automatically and without human interaction 10 minutes after each event, to simulate how ElarmS performs without human assistance, as it would in a real-time application. Eight months of non-interactive operation of the ElarmS simulator have shown that the ElarmS methodology can reliably deliver accurate earthquake

information within a few to a few tens of seconds of event origin. 75 events were successfully processed non-interactively in that time frame. We define an “alarm time” of 4 seconds of data at 4 stations, and at this alarm time the $1-\sigma$ magnitude errors for the non-interactive processing are half a magnitude unit. This is as expected given past performance of the methodology on calibration events and in southern California. At alarm time the ground motion predictions have a $1-\sigma$ error within approximately a factor of 4 in PGA and PGV, or within 0.1 MMI of the actual observed MMI at seismic stations in the BDSN and NCSN networks. These errors are valid for the 75 events in the non-interactive dataset, and it is difficult to estimate the errors for events larger than those in that dataset. The alarm time occurs an average of 15 seconds after event origin. Of the 75 events processed since February 2006, 66 achieve the alarm condition. For these 66 events, ElarmS provides a median warning time of 49 seconds in major Bay Area metropolitan centers.

Two events since February 2006 represent hazardous scenario earthquakes for the Bay Area. In both cases, the magnitude estimate at alarm time is within 0.5 of the network-determined M_L of 4.7, and ground motion predictions are within 0.3 MMI units of the actual peak ground motions observed for the events. Warning times in the Bay Area achieved by the system for the two scenario earthquakes range from about 3 to 30 seconds, depending on the location of the epicenter relative to the city in question.

A functional Earthquake Early Warning system in northern California has the potential to save both lives and money in the event of a major earthquake. Based on the results of simulating the operation of ElarmS, we find the methodology in a condition in which we can move forward to a real-time, online implementation of ElarmS in northern California in the near future.

2.7. Acknowledgments

We wish to thank Doug Neuhauser and Bob Uhrhammer for discussions relating to station equipment and networks, and Dave Wald for discussions about ShakeMap. This paper is Berkeley Seismological Laboratory Contribution #07-05. This work was funded by USGS/NEHRP Grant # 05HQGR0074.

2.8. Figures

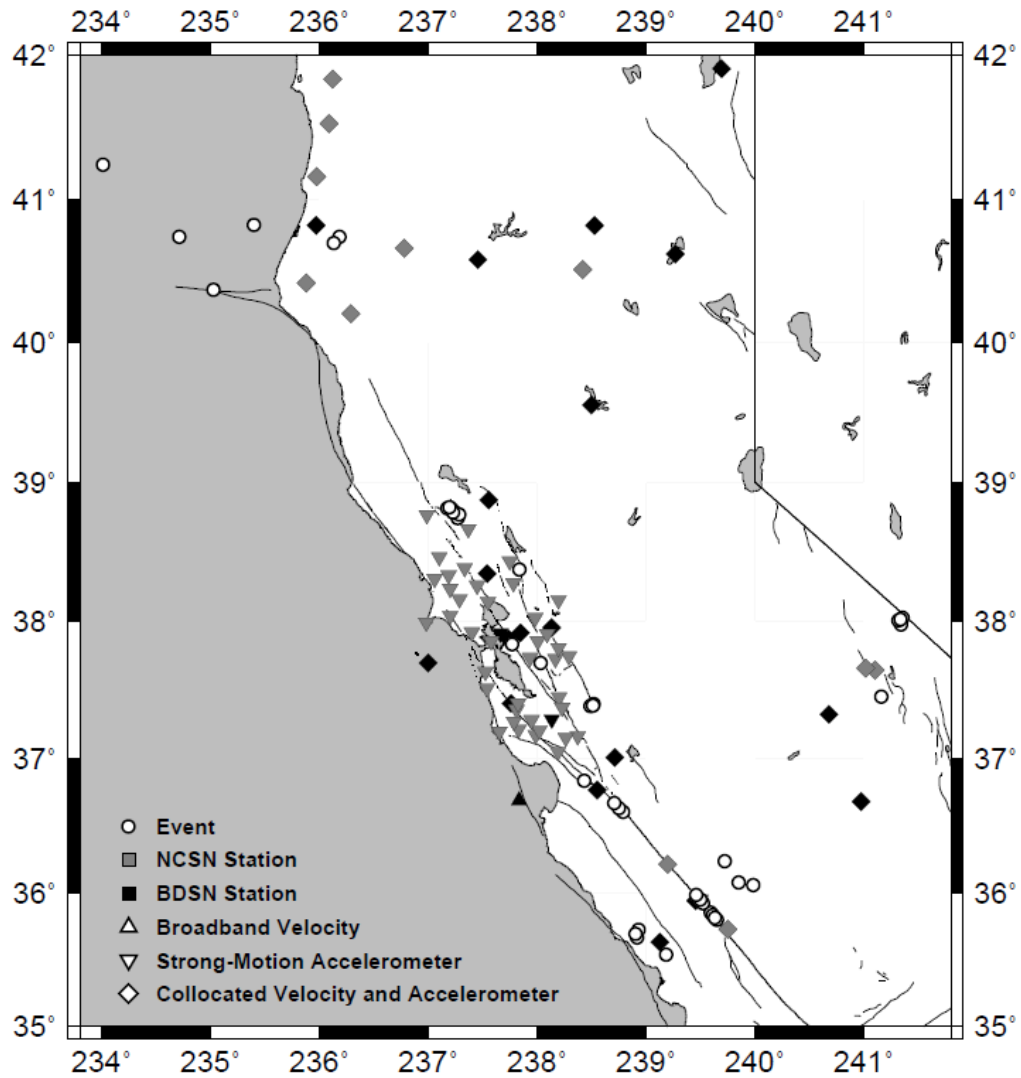


Figure 2.1: Map of California showing distribution of events used in the calibration process (white circles) and stations in the NCSN (gray) and BDSN (black) networks. Upright triangles signify high-gain, broadband velocity sensors. Inverted triangles signify low-gain strong-motion accelerometers, and diamonds signify a station with collocated velocity sensor and accelerometer.

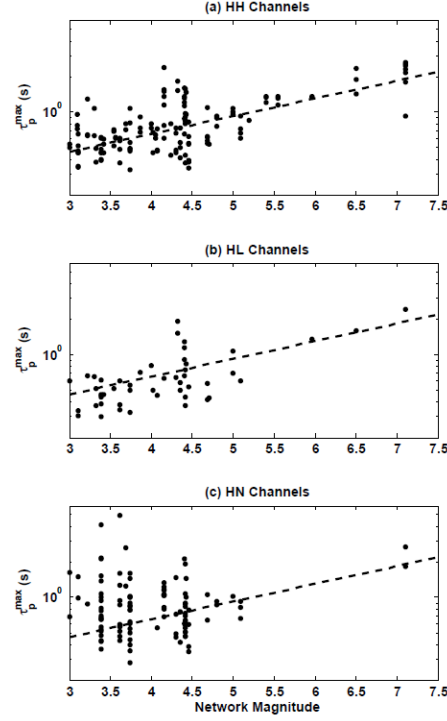


Figure 2.2: Plots of τ_p^{max} vs. magnitude for all calibration events. Each point represents a single station measurement. Measurements are separated by channel code: HH (a) represents velocity sensors, while HL (b) and HN (c) represent accelerometers. The dotted line in all three plots is the same line of best fit using HH and HL data simultaneously.

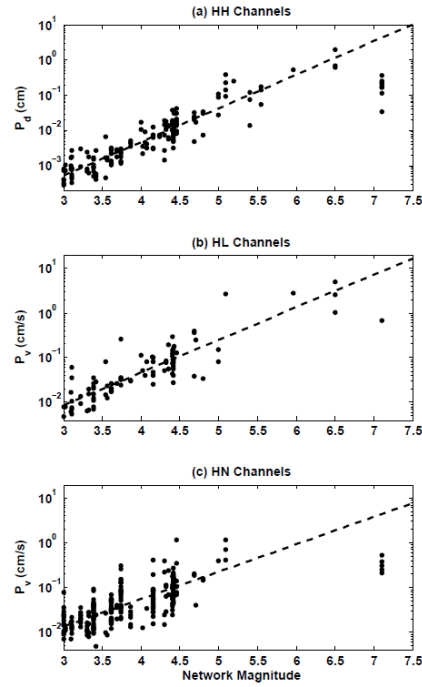


Figure 2.3: Plots of $P_{d/v}$ vs. magnitude for all calibration events. Each point represents a single station measurement, corrected to a distance of 10 km using the empirical best fit equation. Measurements are separated by channel code: HH (a), HL (b) and HN (c). The dotted line in each plot is the line of best fit using data from the corresponding channel only.

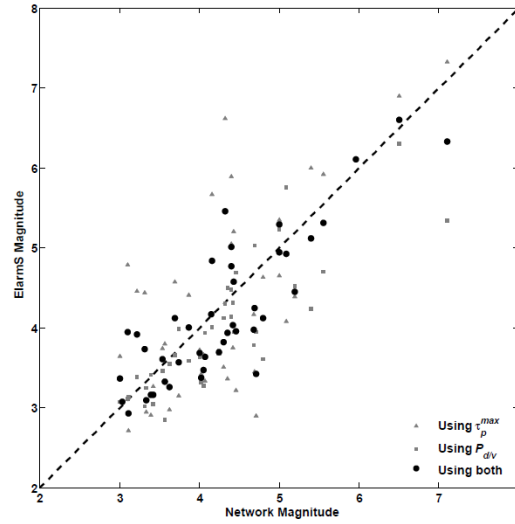


Figure 2.4: Plot of ElarmS magnitude estimate vs. network-derived magnitude (usually M_L or M_w) for the calibration events. Gray triangles are magnitude estimates using τ_p^{max} only, gray squares are estimates using $P_{d/v}$ only, and black circles are a linear average of the two magnitudes for each event.

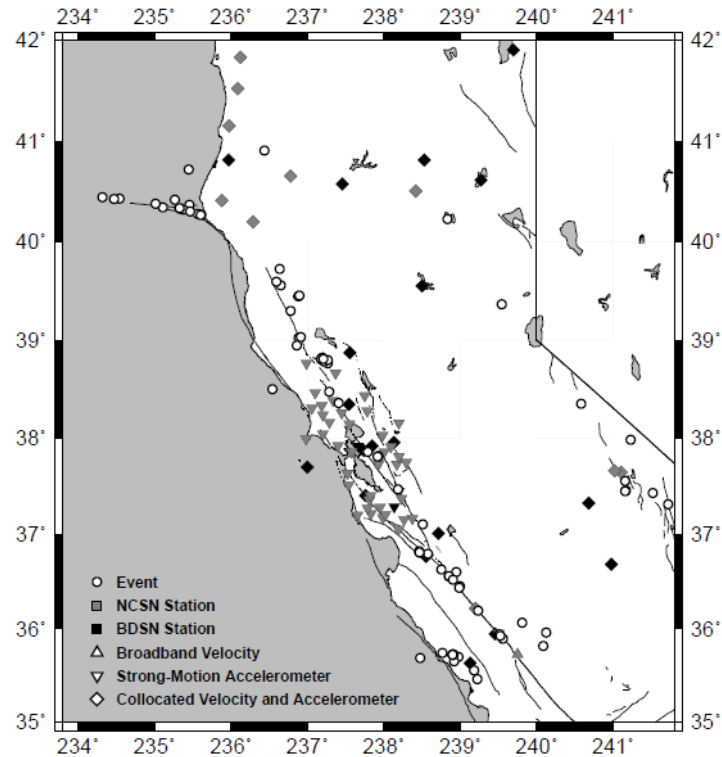


Figure 2.5: Map of California showing distribution of events (white circles) processed non-interactively by ElarmS from February through September 2006. Stations in the NCSN (gray) and BDSN (black) networks are plotted as upright triangles for high-gain, broadband velocity sensors. Inverted triangles signify low-gain strong-motion accelerometers, and diamonds signify a station with collocated velocity sensor and accelerometer.

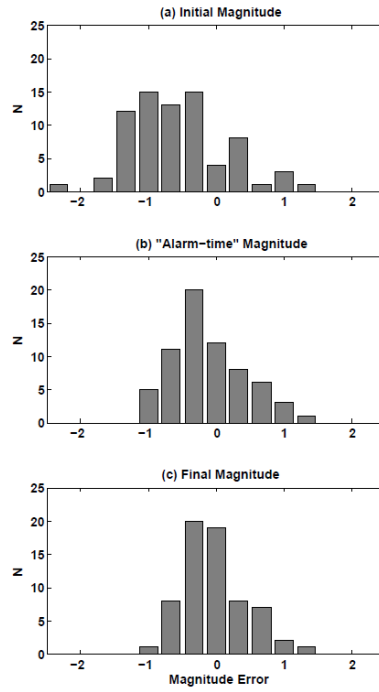


Figure 2.6: Histograms showing errors in magnitude estimate for all non-interactive events (a) at one second after detection (75 events), (b) at “alarm time” when 4 seconds of data are available in 4 channels (66 events), and (c) 60 seconds after the event began for those events which achieved the alarm condition (66 events).

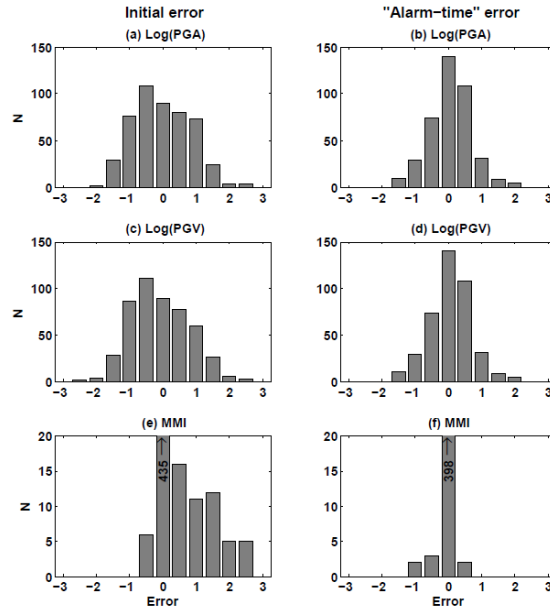


Figure 2.7: Histograms showing ground motion prediction errors at all stations for all non-interactive events. Only stations which had not already observed peak motions at the initial and alarm times are included in the data. (a) Errors in Log(PGA) at one second after detection, and (b) at “alarm time” when 4 seconds of data are available in 4 channels; (c) errors in Log(PGV) at one second after detection and (d) at alarm time; (e) errors in MMI at one second after detection and (f) at alarm time. In (e) and (f), the 0-centered bin is off scale. The value of this bin is 435 in (e) and 398 in (f).

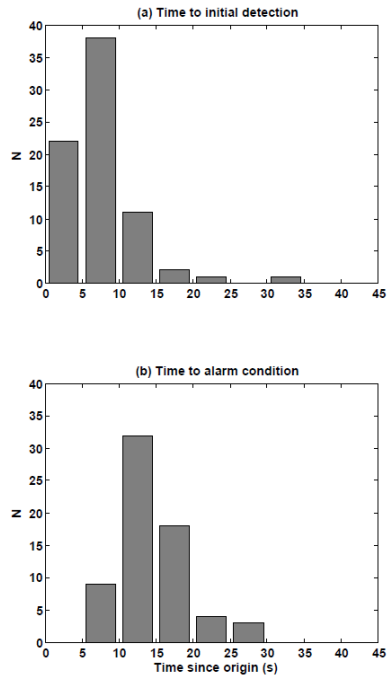


Figure 2.8: Histograms of time between event origin and initial detection (a) and between event origin and the “alarm time” at which 4 seconds of data are available in 4 channels (b) for all non-interactive events.

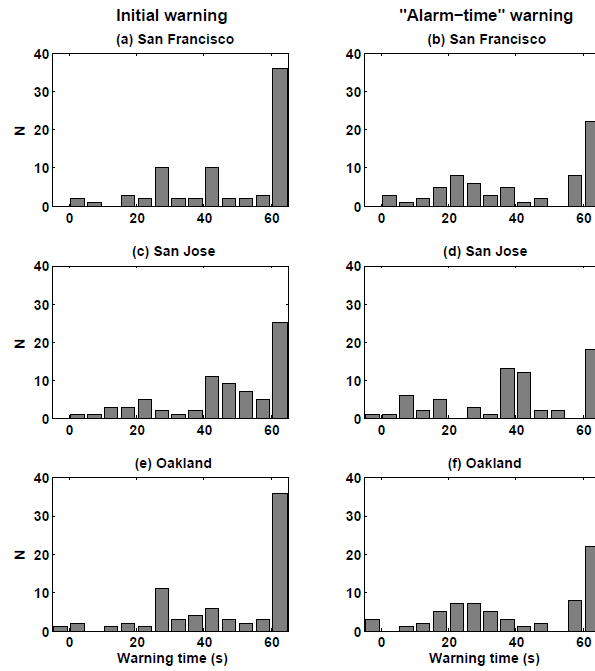


Figure 2.9: Histograms of warning time until onset of largest ground motions at San Francisco, San Jose and Oakland for all non-interactive events. Onset times are estimated using a move-out speed of 3.75 km/s. Warning times are calculated from one second after initial detection (a, c and e) and from “alarm time” (b, d and f), defined as having 4 seconds of data in 4 channels.

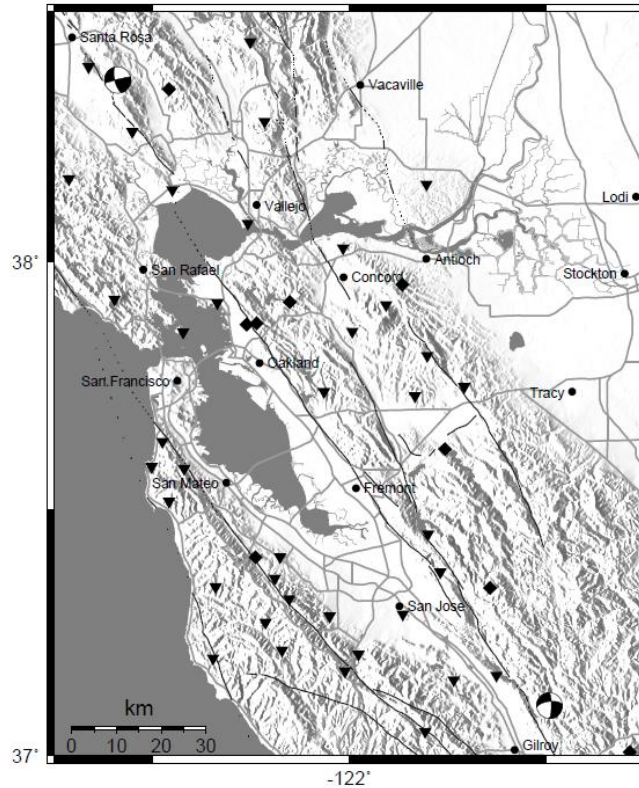


Figure 2.10: Map of the San Francisco Bay Area showing the location and focal mechanisms of two hazardous scenario earthquakes processed non-interactively by ElarmS. The black symbols are stations in the NCSN and BDSN networks. Stations are plotted as upright triangles for high-gain, broadband velocity sensors, inverted triangles for low-gain strong-motion accelerometers, and diamonds for stations with collocated velocity sensor and accelerometer. Focal mechanisms are from the BDSN regional moment tensor catalog [Pasyanos et al., 1996].

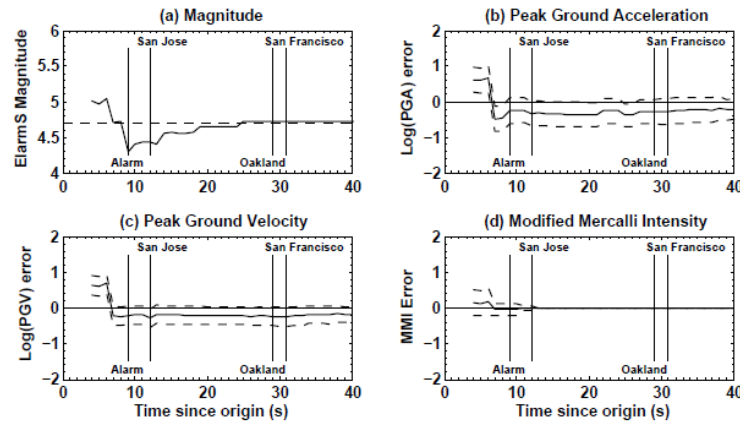


Figure 2.11: Plot of ElarmS output for the first 40 seconds of the Gilroy event. Vertical lines represent (as labeled) the time of alarm condition, or the onset of largest ground motions at San Francisco, San Jose or Oakland based on a move-out of 3.75 km/s. (a) Magnitude estimate. The dotted horizontal line represents the network-based M_L of 4.7. (b) Error in the logarithm of predicted peak ground acceleration. At each time interval, the plot incorporates all stations which have not yet observed peak ground motions. Dotted lines represent the $1-\sigma$ error envelope. (c) Error in the logarithm of predicted peak ground velocity. (d) Error in the predicted Modified Mercalli Intensity.

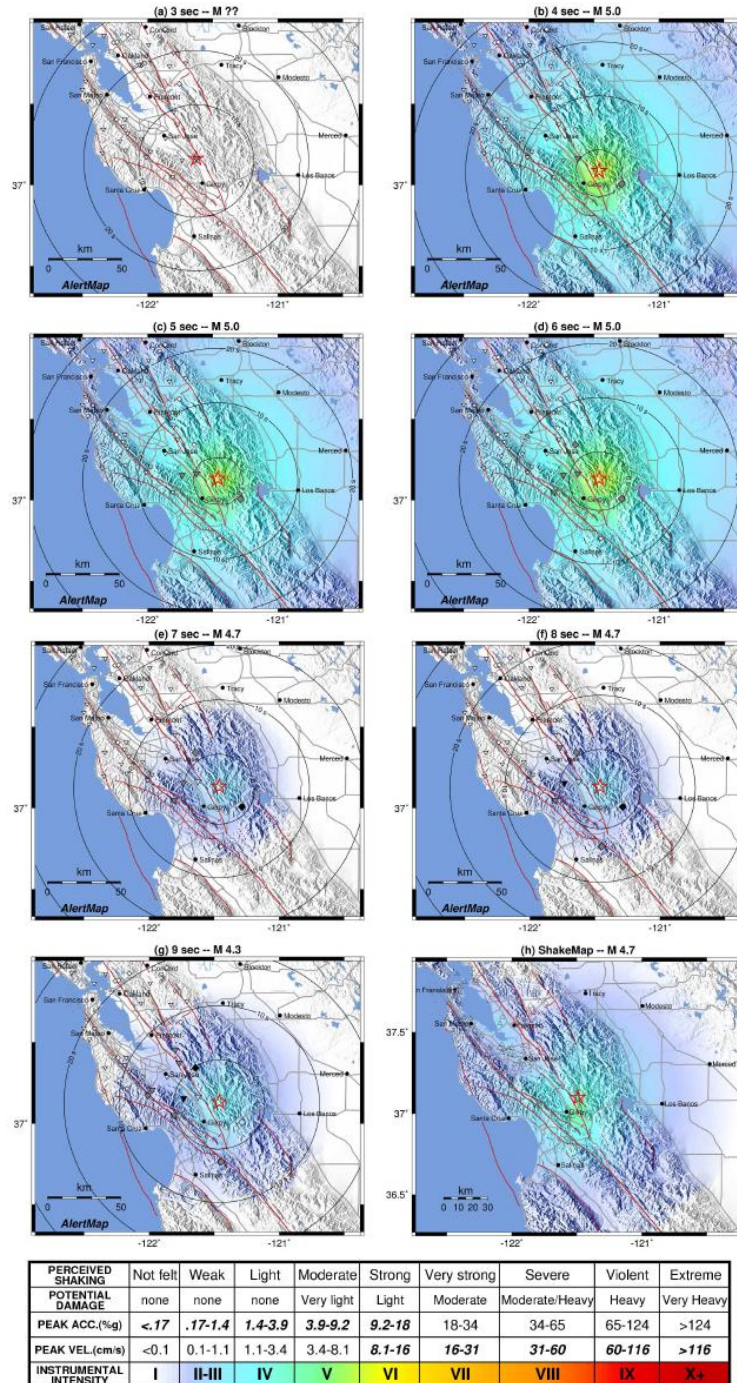


Figure 2.12: (a-g) ElarmS AlertMap output for 3-9 seconds after the origin of the Gilroy event. Time since event origin and the magnitude estimate are shown above each AlertMap. The epicenter is plotted as a red star. Stations which have triggered are shown in gray, stations which are experiencing peak ground motions are black, and those that have reported peak ground motion are color coded according to the scale at bottom. Stations are plotted as upright triangles for high-gain, broadband velocity sensors, inverted triangles for low-gain strong-motion accelerometers, and diamonds for stations with collocated velocity sensor and accelerometer. The circular contours represent time until onset of strong ground motion based on the location and origin time of the event and a move-out of 3.75 km/s. The color field is the ElarmS prediction of Modified Mercalli Intensity, according to the scale at bottom. (h) The ShakeMap for the Gilroy event. The color field represents actual instrumentally-observed Modified Mercalli Intensity for the event, processed after the event occurred.

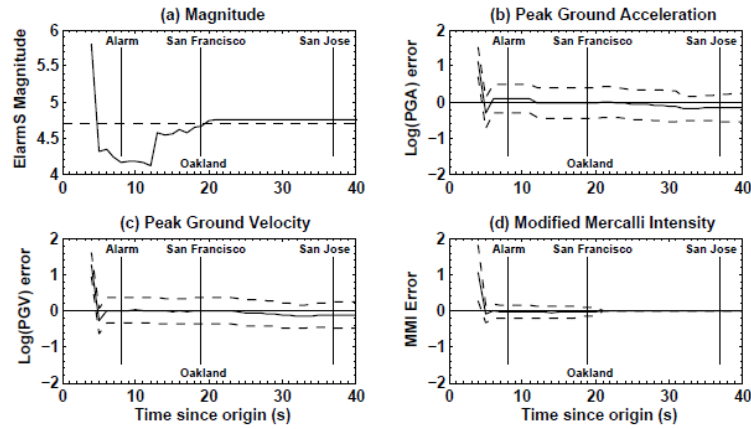
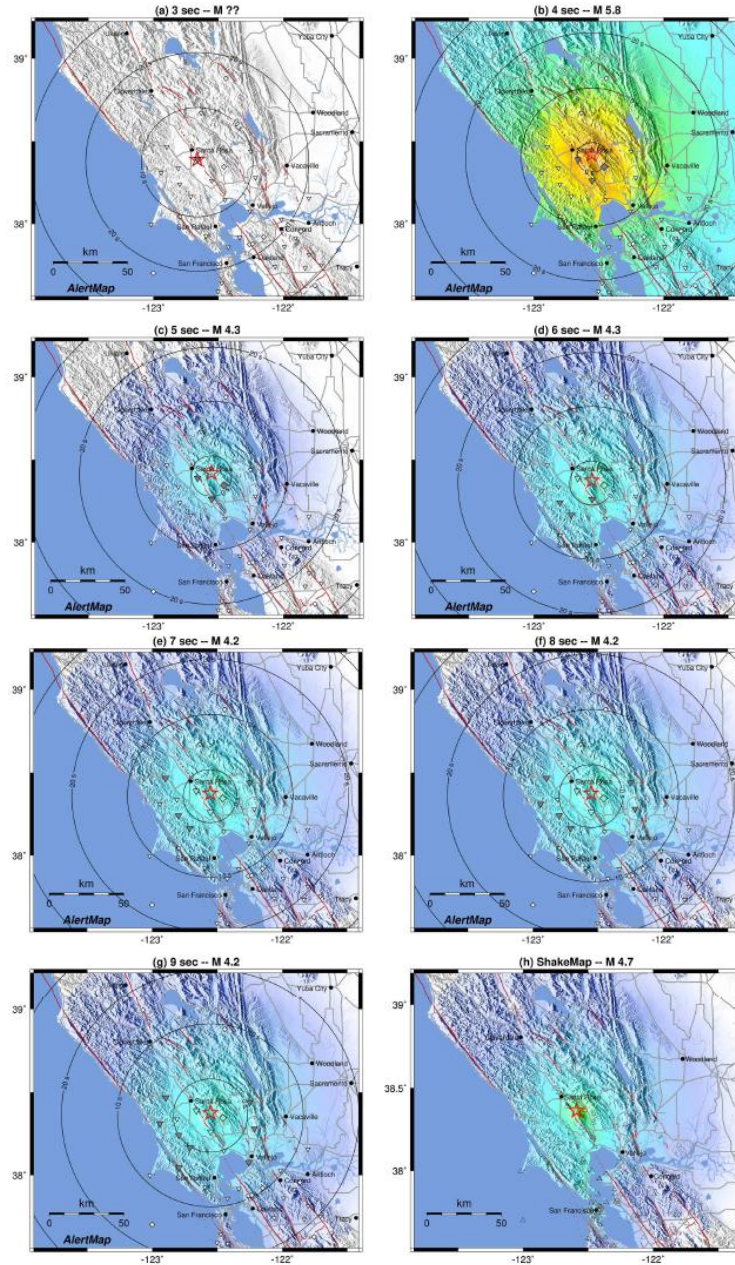


Figure 2.13: Plot of ElarmS output for the first 40 seconds of the Santa Rosa event. Vertical lines represent (as labeled) the time of alarm condition, or the onset of largest ground motions at San Francisco, San Jose or Oakland based on a move-out of 3.75 km/s. (a) Magnitude estimate. The dotted horizontal line represents the network-based M_L of 4.7. (b) Error in the logarithm of predicted peak ground acceleration. At each time interval, the plot incorporates all stations which have not yet observed peak ground motions. Dotted lines represent the $1-\sigma$ error envelope. (c) Error in the logarithm of predicted peak ground velocity. (d) Error in the predicted Modified Mercalli Intensity.



PERCEIVED SHAKING	Not felt	Weak	Light	Moderate	Strong	Very strong	Severe	Violent	Extreme
POTENTIAL DAMAGE	none	none	none	Very light	Light	Moderate	Moderate/Heavy	Heavy	Very Heavy
PEAK ACC.(%g)	<.17	.17-1.4	1.4-3.9	3.9-9.2	9.2-18	18-34	34-65	65-124	>124
PEAK VEL.(cm/s)	<0.1	0.1-1.1	1.1-3.4	3.4-8.1	8.1-16	16-31	31-60	60-116	>116
INSTRUMENTAL INTENSITY	I	II-III	IV	V	VI	VII	VIII	IX	X+

Figure 2.14: (a-g) ElarmS AlertMap output for 3-9 seconds after the origin of the Santa Rosa event. Time since event origin and the magnitude estimate are shown above each AlertMap. The epicenter is plotted as a red star. Stations which have triggered are shown in gray, stations which are experiencing peak ground motions are black, and those that have reported peak ground motion are color coded according to the scale at bottom. Stations are plotted as upright triangles for high-gain, broadband velocity sensors, inverted triangles for low-gain strong-motion accelerometers, and diamonds for stations with collocated velocity sensor and accelerometer. The circular contours represent time until onset of strong ground motion based on the location and origin time of the event and a move-out of 3.75 km/s. The color field is the ElarmS prediction of Modified Mercalli Intensity, according to the scale at bottom. (h) The ShakeMap for the Santa Rosa event. The color field represents actual instrumentally-observed Modified Mercalli Intensity for the event, processed after the event occurred.

3. Statistical Testing of Theoretical Rupture Models Against Kinematic Inversions

Gilead Wurman, Richard M. Allen, Douglas S. Dreger and David D. Oglesby

Submitted to the *Bulletin of the Seismological Society of America* as:

Wurman, G., R. M. Allen, D. S. Dreger, and D. D. Oglesby (submitted 12 March 2010), Statistical testing of theoretical rupture models against kinematic inversions.

3.1. Abstract

A central question in earthquake physics is whether earthquakes are a fundamentally self-similar process or not. This issue continues to be vigorously debated in the seismological community. If they are not self-similar phenomena, small ruptures may behave qualitatively differently from large ruptures. In particular, this may be observable in the evolution of slip over the course of an earthquake. Modeling studies have attempted to answer this question, but there exist models that show self-similar behavior and others that show deterministic behavior.

We examine 167 kinematic slip inversions, most from the SRCMOD database (see Data and Resources Section), to determine whether these models exhibit self-similar characteristics or scaling with magnitude. The large number of models available for the analysis allow us to address statistically the early slip history which otherwise, taken for each model individually, is not sufficiently well-resolved to draw robust conclusions.

We find a clear and robust scaling of slip amplitude within the first second to the first 10 seconds of rupture with the final size of the earthquake. We separate out the models according to use of linear or nonlinear methods and find no distinguishable difference between different types of inversions. The scaling becomes more pronounced when earthquakes with large grid spacing are removed from the analysis, suggesting that the coarser parameterization of larger events is not responsible for the observed behavior. This scaling is suggestive of a degree of determinism in earthquake rupture, though it is likely that we are observing the aggregate behavior of individually cascading ruptures, and not the fundamental physics of each rupture.

3.2. Introduction

The process by which earthquake ruptures initiate and propagate is usually expressed as one of two broadly-defined mechanisms. The cascade model [*Bak and Tang, 1989; Steacy and McCloskey, 1998; Sato and Kanamori, 1999; Ide and Aochi, 2005; Otsuki and Dilorv, 2005*] explains earthquake rupture as a process by which a small initial patch of the fault nucleates slip, and the stress concentration at the edges of that patch load the adjacent patch which may or may not rupture as a result of the stress from the first patch. If it ruptures, it subsequently loads yet another patch, which in turn may fail as a result. The final size of the earthquake is determined by whether each subsequent patch does or does not fail from the loading of the previous patch. A consequence of this behavior is that the initial rupture of a large earthquake is indistinguishable from that of a small earthquake; it is the spatial distribution of stress (i.e., the presence or absence of barriers and asperities) that determines whether the final earthquake will be small or large. Thus, this model predicts that earthquakes are self-similar phenomena, and are therefore inherently non-deterministic

in that one cannot predict the final size of the rupture while the event is ongoing. The second model, known as the preslip model [*Ellsworth and Beroza*, 1995; *Aki*, 2000; *Ohnaka*, 2000; *Ohnaka*, 2004; *Abercrombie*, 2005], suggests that a nucleation slip patch, which is either aseismic or seismically slow, is allowed to develop before the fast rupture process begins. This patch loads the surrounding fault according to the amount of pre-seismic slip accumulated, and this variation accounts for any difference between large and small earthquakes. The consequence of this model is that the rupture of a large earthquake is immediately distinguishable from the rupture of a small earthquake because of differences in stress in the nucleation region. Thus the final size of the event is deterministic, having been set by the amount of preslip before the onset of seismic rupture.

There is a diversity of modeling results that alternately support either cascade- or preslip-type rupture. These include both data-driven kinematic inversions of real earthquakes and more theoretical, dynamic forward models that start with assumptions on the initial fault stress and constitutive properties, and calculate the evolution of fault rupture and slip. For example, *Ohnaka* [2000, 2004] developed constitutive models based on fracture mechanics, in which the nucleation phase scales with the size of the event, whereas *Lapusta and Rice* [2003] created simulations in which the aseismic nucleation phase did not scale with the size of the final event. *Aochi and Ide* [2004] found in their numerical models that a scale-dependent nucleation phase is not necessary to generate a range of earthquake sizes, but *Fukuyama and Madariaga* [2000] and *Fukuyama et al.* [2002] developed models in which a nucleation phase was observed, similar to the observations of *Ellsworth and Beroza* [1995, 1998]. *Mai et al.* [2005] examined the correlation between the hypocenter location and the region of peak slip in around 90 kinematic slip inversions of over 50 global earthquakes between M_w 4.1 and 8.1. They found that the focus of many large earthquakes is not collocated with the point of peak slip. While this result is suggestive of cascade-type behavior, they determined that the focus is statistically more likely to fall near a region of high slip than low slip. This observation may be an indication that a focus near a high-stress region generates a more energetic initial rupture and tends to generate large earthquakes. *Oglesby and Day* [2002] indirectly address this issue with dynamic models, noting that with rough initial stress distributions only earthquakes that nucleate near relatively large regions of high average shear stress can develop into full spontaneous ruptures. *Ripperger et al.* [2007] find similar results in their simulations of ruptures with heterogeneous initial stress conditions, though this might be due in part to some ruptures not having an initial patch size large enough to generate a self-propagating rupture. *Murphy and Nielsen* [2009] find, using dynamic models, that altering the initial stress conditions on the fault plane can

control the evolution of rupture, even when the nucleation conditions do not vary between simulations.

At least some of the disagreement between the findings of the data-driven studies may be attributed to the high degree of variability among kinematic slip inversions, even for the same event [Beresnev, 2003]. The recent development of the SRCMOD database (see Data and Resources Section), a large database of kinematic slip inversions, makes it possible to examine many fault models in a statistical fashion to suppress the effects of this variability. We examine the aggregate behavior of faults in many different earthquakes and slip models to characterize the evolution of slip in the early stages of earthquake rupture. Although the slip at the beginning of rupture is difficult to resolve in kinematic inversions, using such a large number of events allows us to make first-order observations of any relationships that might be present. By the nature of this analysis, it is not possible to distinguish whether any relationships result from behavior of individual ruptures or from aggregate behavior of the sample population as a whole. The purpose of this work, then, is to assess whether earthquakes as a statistical population, viewed through kinematic rupture models, exhibit any deterministic or self-similar behavior, rather than to infer the fundamental physics of earthquake rupture from a statistical treatment with large inherent uncertainties.

3.3. Method

We examine 152 inversions of 80 different earthquakes in the SRCMOD database (see Data and Resources Section) as well as 7 teleseismic and 8 joint geodetic/teleseismic inversions provided by M. E. Pritchard [Pritchard et al., 2006; Pritchard et al., 2007; Pritchard and Fielding, 2008; Loveless et al., in review] for this study, for a total of 167 inversions and 95 events (Figure 3.1). The events range in magnitude from 4.1 to 8.9 and include a variety of mechanisms in a variety of tectonic regimes.

We take the final slip distribution for each inversion and reconstruct the time-evolution of slip on the fault from rupture time and rise time information. We calculate the moment release within a given time window by summing the moment at each grid point via the relation $M_0 = \mu SD$ where

$\mu = 3 \times 10^{11}$ dyne/cm², S is the area of each grid element and D is the slip at each grid element. 53 models have point-by-point rupture time data and 24 of these also have point-by-point rise time data. One model has only rise-time data with no rupture time information. In cases where point-wise data is available we initiate slip on each grid point at the associated rupture time, and increase slip to the final amount in a linear ramp over the associated rise time [Haskell, 1964]. In models where one or both of these parameters is not recorded point-wise we take the reported average rupture velocity or

rise time and assume a rupture front expanding isotropically from the hypocenter. 24 models have neither point-wise rupture time information nor reported average rupture velocity, and for those events we assume a rupture velocity of 3.0 km/s (the most common average rupture velocity in the database). There are also 15 events with no point-wise rise time information and no reported average rise time, and for these events we assume the slip rises instantaneously. This is not a realistic assumption, but only 15 events are subject to this assumption and the overall results are invariant under assignment of rise times up to 10 seconds; thus we choose zero rise time for simplicity.

3.4. Relationship between early and final moment

Before inspecting the slip models for any correlation we must consider the expected outcomes in terms of the competing end-member hypotheses of a preslip (deterministic) model and a cascading (self-similar) model. The latter end-member case implies that at any given time after rupture initiation all earthquakes look the same. That is, 1 second into the rupture process all earthquakes should, on average, have the same magnitude as an earthquake with duration of 1 second. To first order, we can approximate this magnitude via the relation

$$M_0 = \mu SD = \mu \cdot \pi(V_R t)^2 \cdot \alpha V_R t = \pi \mu \alpha (V_R t)^3 \quad (3.1)$$

where $\mu = 3 \times 10^{11}$ dyne/cm², $S = \pi(V_R t)^2$ is the fault area after t seconds assuming a rupture velocity $V_R = 3 \times 10^5$ cm/s, and $D = \alpha V_R t$ is the mean slip with $\alpha = 10^{-5}$ [Heaton, 1990; Wells and Coppersmith, 1994]. The choice of the value of α is driven by the kinematic data themselves, and is discussed at length in the following section. Using this approximation, we find that a source duration of 1 second corresponds to a moment of approximately 2.5×10^{23} dyne-cm, or M_w 4.9. Thus in the cascading end-member case, any earthquake larger than magnitude 5 should look like a magnitude 5 at one second after nucleation. We can visualize this hypothesis in Figure 3.2. Allowing for a normally distributed variability around the theoretical magnitude at 1 second, the data should resemble the randomly generated dataset shown in Figure 3.2, and have a best-fit slope around zero. The solid diagonal line in Figure 3.2 represents the limit in which the initial magnitude after 1 second is equal to the final magnitude, meaning the rupture has propagated to completion. Since there is no way for the initial magnitude to exceed the final magnitude, the triangles in Figure 3.2 are physically impossible and in reality no points will lie above this solid line. This can

potentially introduce a spurious positive slope to the data, which we assess in greater detail in the next section. We minimize this potential bias in the following analyses by culling all data points for which the final magnitude is less than the reference magnitude (represented as hollow symbols in Figure 3.2), i.e., points that lie to the left of the intersection between the solid diagonal line and the dotted horizontal line. The hypothesis for a deterministic model is that there will be some positive scaling of the magnitude at 1 second with the final magnitude of the earthquake, yielding a positive slope to the data. The cascading model will be the null hypothesis against which we will test the data, using a one-sided t-test to determine whether the best-fit slope is statistically greater than zero. Because the data have a large variance and we are trying to resolve early slip which is inherently difficult, we will require a 95% confidence level to reject the cascading model hypothesis.

We first examine the dataset in as complete a form as possible, to gain the greatest possible statistical advantage from the large number of data points. Figure 3.3 shows the initial magnitude plotted against the final magnitude for each event (as determined by summing the completed slip over the entire fault plane) for four time windows ranging from 1 second to 8 seconds. At this stage we apply only the most rudimentary checks on the data to avoid introducing biases into the analysis. We manually pick outliers (shown as crosses on all figures), and exclude any events for which the initial magnitude is 99% or more of final magnitude, as those events have effectively terminated by the end of the time window. We also exclude all events with a final magnitude less than the reference magnitude for the null hypothesis, because these events are expected to have a duration less than the time window. Theoretically, had they continued rupturing until the end of the time window, they would have an initial magnitude greater than their final magnitude. Since this is not possible (there cannot be any points above the solid line in the figures) these events can only introduce a bias of the best-fit line to greater slope, unrealistically favoring deterministic behavior.

In Figure 3.3(a) through (d) the null hypothesis can be rejected with greater than 95% confidence. In this figure as well as Figure 3.4 circles represent inversions with no point-wise rupture time data, triangles represent inversions with point-wise rupture time information, and hollow symbols represent data for which the time window is shorter than the rise time, or for which only one grid point has begun rupturing within the time window. The slope of the best-fit lines for 1, 2, 4 and 8 second time windows are strongly positive, suggesting some degree of non-self-similar behavior for these models. This result holds for all time windows between 1 and 10 seconds. For time windows between 8 and 10 seconds the confidence, while still exceeding the 95% level, is not as high as for time windows between 1 and 7 seconds (Figure 3.5). We therefore expect the initial magnitude to scale more strongly with final magnitude for longer time windows. One

possible explanation for the degradation in scaling for longer time windows is that more and more events are being excluded due to having completed rupture, thus reducing the number of data points available for analysis. In Figure 3.3(a-c) the number of points used for the fit varies between 112 and 128, and by 8 seconds (Figure 3.3d) that number has fallen to 80. Another possibility is that longer time windows afford greater resolution of the slip within the time window, implying that the strong correlation observed for shorter time windows is a spurious result of poorly resolved slip in such short time spans. The influence of poorly resolved slip can be approximated visually by noting the open symbols, which represent models for which the time window was either shorter than the average rise time for the model or for which only one grid element had begun slipping in that time window.

In Figure 3.4 we attempt to reduce the influence of poorly resolved slip in two ways. Primarily, we disregard all of the “open” data points from Figure 3.3 which represent cases where the slip is likely to be particularly poorly resolved owing to the time window being too short. In addition, we recalculate both the initial and final magnitude for each point, disregarding any slip which is less than 10% of the peak slip for the model. This process is to account for the fact that slip below 10% of peak slip is generally regarded as being poorly resolved in kinematic inversions and thus an unstable component of the slip models [Somerville et al., 1999]. Remarkably, the correlation between early and final magnitude is now even stronger, with the null hypothesis being rejected at greater than 95% confidence for all time windows as seen in Figure 3.5. This suggests that poor resolution of slip in short time windows is not generating a spurious correlation between early and final magnitude. Rather, the analysis suggests that the decreasing number of data points in longer time windows (owing to more ruptures having run to completion) is primarily responsible for the weaker correlation for 8-10 second time windows seen in Figure 3.3 and Figure 3.5.

There are two additional features of both Figure 3.3 and Figure 3.4 that argue against purely cascading behavior in these slip inversions. The first is that the large majority of points in all time windows lie above the null hypothesis line (dotted lines in Figure 3.3 and Figure 3.4), indicating that these models as a rule exhibit greater than expected moment release at any given point in the rupture. Recalling the null hypothesis (Figure 3.2), the cascade model predicts that moment release after a given time window will be normally distributed above and below the null hypothesis line. That is, some models will exhibit greater than expected early moment release, and others will exhibit less than expected moment release at a given point in time. Since this is quite clearly not the case with these data, the implication is that the events examined here do not obey the cascade model. The other feature is that the best-fit lines all intercept the 1:1 line on the plots close to the null hypothesis line. This means that whatever scaling is seen converges

to the correct magnitude, i.e. a 1 second rupture for a magnitude 5 earthquake releases about a magnitude 5, as it should. While neither of these observations is necessarily robust on its own, they serve as corroboration that the observed scaling is physically realistic.

3.5. Effects of Size Distribution, Bias and Stress Drop

3.5.1. Size Distribution

One possible confounding effect in the foregoing analysis is the uneven distribution of event sizes in the dataset. This distribution is atypical of earthquake catalogs, in that the number of events is limited both on the high end due to the small number of events at large magnitudes, and at the low end due to the dearth of small events which are well-recorded enough to be modeled. The magnitude distribution in the dataset is approximately normal with a mean magnitude of 6.75 ± 0.65 . The potential control exerted on the best-fit line by the few events at large magnitude is thus balanced by the few small-magnitude events. To assess the degree of this effect and any bias it introduces we randomly select three events in each bin from magnitude 4.0 to 9.0 in $\frac{1}{4}$ magnitude increments and perform the analysis on this subset of data which is approximately homogeneously distributed across the entire magnitude range. The results of 100 such random subsets are presented in Figure 3.6. This figure shows that homogenizing the size distribution tends to reduce the confidence with which the null hypothesis can be rejected, but that this degradation of confidence is more pronounced for shorter time windows. This result is to be expected as the subsampling is more severe for the shorter time windows that have more data to begin with. The net effect is to even out the confidence levels between the different time windows, so that we no longer observe the degradation in confidence at longer time windows as compared with the shorter windows. This result again suggests that the degradation observed in the previous section is primarily due to a reduced number of data points at the longer time windows. In any case, all but one of the time windows still show the null hypothesis rejected at 95% confidence, leading us to conclude that the heterogeneous distribution of event sizes is not a controlling factor in our prior results.

3.5.2. Bias

As previously stated and shown in Figure 3.2, the fact that no data can lie in the upper left half of the plane introduces an artificial upward bias to the best-fit slope in the foregoing analyses. We attempt to minimize the effect of this bias by culling all data for which the final magnitude is less than the theoretical reference magnitude (open symbols in Figure 3.2), but some bias will remain because of points above the 1:1 line in Figure 3.2 (solid triangles) which have a magnitude greater than the reference magnitude,

but are nevertheless physically impossible. The possibility exists that our conclusions thus far have been controlled by this unavoidable bias. Assuming a normally distributed random error about the reference magnitude as in Figure 3.2, the degree of bias is controlled by the standard deviation of that error and by the reference magnitude, which affects how many of the low magnitude data are culled. To estimate the effect of this bias, we generate a set of 200 points randomly normally distributed about the reference magnitude of 6.9 (corresponding to a 10-second time window) with a standard deviation of 0.38 magnitude units. This standard deviation was chosen by calculating the standard deviation of the unfiltered dataset in Figure 3.3 for all time windows between 1 and 10 seconds. We calculate a one-sided standard deviation considering only those points falling below the best-fit magnitude, to avoid artificially reducing the standard deviation due to the physical limit of the 1:1 line. Using this technique a standard deviation of 0.38 magnitude units, corresponding to the 2-second time window, is the greatest standard deviation over all the time windows. We cull the random dataset as described above and in Figure 3.2, and take the best-fit slope of these data. We do this 1000 times and take the average of the best-fit slope over all simulations and time windows, a slope of 0.105 for a 10 second time window, to be the biased null hypothesis. Since we only simulate points up to magnitude 9, by using the reference magnitude of 6.9 for the 10 second time window we limit the possible magnitudes to be between magnitudes 7 and 9, thus maximizing the potential bias. We can test the data, as in Figure 3.5, against this biased best-fit slope rather than against a slope of zero. Figure 3.7 shows the probability of acceptance of this new null hypothesis as a function of time window length. For all time windows less than 8 seconds, the biased null hypothesis of a slope less than or equal to 0.105 is still rejected with greater than 95% confidence, even in the unfiltered dataset. We therefore conclude that the artificial bias introduced by our method is not controlling the observed scaling relations.

3.5.3. Stress drop

The reference magnitude is itself affected by the choice of the parameter α , which is related to stress drop via the relation $\Delta\sigma = \mu D/L = \mu \cdot \alpha$. Our choice of $\alpha = 10^{-5}$ represents a stress drop of 3×10^6 dyne/cm², or 3 bars. Although this is a very low stress drop, the choice of $\alpha = 10^{-5}$ is driven by the dataset itself, as shown in Figure 3.8. This figure shows the actual magnitude of each inversion calculated by summing all slip greater than 10% of peak over the fault plane, plotted against the duration of the ruptures. Using only the slip greater than this threshold yields better-constrained magnitudes [Somerville et al., 1999], though incorporating these small slips adds less than 1.3% to the final magnitude of any event. The duration of each rupture is calculated by taking all the grid-point rupture times for points with slip greater than 10% of peak. Simply taking the maximum time at which a grid

point with sufficient slip ruptured would bias the analysis toward excessively long rupture times, leading to low apparent stress drops. Instead, we take the total rupture duration to be the 90th percentile rupture time, meaning that 90% of points with sufficient slip have ruptured prior to this time and 10% will rupture after this time. The lines superimposed on this data represent magnitude vs. duration calculated via Eq. 3.1 for stress drops of 0.3, 3, 30 and 300 bars (α from 10^{-6} to 10^{-3}). Figure 3.8 clearly shows that a stress drop of 3 bars is more consistent with the data than the canonical 30 bars. While the finite-source models in the database tend to have low average stress drop when calculated in this way, if we consider the stress drops only over primary asperities rather than overall ruptures they would be larger due to the non-uniform nature of the slip. On the other hand, to make the rupture durations fit the 30 bar line equally well as the 3 bar line, we would need to set the rupture duration to the 60th percentile rupture time rather than the 90th. That is, the analysis would suggest that kinematic inversions tend to overestimate rupture duration by a factor of 5/3, which is highly unlikely to be true.

Although the choice of $\alpha = 10^{-5}$ appears reasonable given the dataset, it is worth examining the effect that a different choice of α has on the foregoing analysis. The more severe culling (reference magnitude increases from 6.9 to 7.6 for a 10 second time window) has a predictably adverse effect on the confidence with which we can reject the null hypothesis. Figure 3.7 shows the likelihood of accepting the null hypothesis of zero slope as a function of time window length. With high quality data it is still possible to reject the null hypothesis with as much as a 6 second time window, but the severe culling makes this impossible to do with unfiltered data. This is because filtering out the poorly-resolved models preferentially removes small-magnitude events (which are much more likely to be poorly resolved as discussed regarding Figure 3.4). Consequently, the more severe low-magnitude cutoff has less of an effect on the high-quality data since most of the low-magnitude events have been filtered out already.

An additional effect of increasing the stress drop is that the biased best-fit slope increases from 0.105 to 0.221, due to the more severe culling of events. If we use this biased slope as the null hypothesis, it becomes very difficult to reject (Figure 3.7). The unfiltered data cannot reject the null hypothesis with any confidence (for longer time windows the best-fit slope is even less than 0.221), and even the high quality data can only reject the null hypothesis for time windows less than 3 seconds. This demonstrates that using a higher average stress drop with this analysis introduces a very severe bias which renders the conclusions unsupportable. Although there is still an observable correlation between early and total moment release, it is indistinguishable from the bias inherent to the method.

3.6. Variability due to inversion data and methods

It is apparent from the foregoing analysis that kinematic slip models favor a larger than expected amount of slip early in the rupture process, which supports the possibility that the aggregate behavior of earthquakes is not entirely self-similar. There is still quite a large variability in the data, which may be due in some part to the diverse dataset used in the foregoing analysis, including both linear and nonlinear inversions as well as models incorporating strong motion, teleseismic, geodetic and other data sources. Of interest is whether one class or another of inverse techniques particularly favors early slip over late slip, or if incorporating a particular type of source data biases the slip distribution in one or the other direction.

To begin with, we break out the models by the type of data they incorporate based on the SRCMOD classification (see Data and Resources Section) which identifies eight data types: strong ground motion (SGM) seismic, teleseismic, trilateration, leveling, GPS, InSAR, surface rupture mapping, and other. The "other" classification primarily refers to tsunami modeling studies using the methods of Satake and Tanioka [*Satake, 1987; Tanioka and Satake, 2001*]. The GPS classification refers mostly to static coseismic displacements, but one model (Ji's 2004 Parkfield model, see Data and Resources Section) includes both static-offset GPS and 1-Hz continuous GPS displacement waveforms, which we treat as strong ground motion records for the purpose of this analysis. Only one model, the 2002 Denali rupture model of *Oglesby et al. [2004]* incorporates surface rupture mapping information in the inversion. For the purpose of the following analysis we separate the eight data types into two classes: "seismic" data which includes strong ground motion and teleseismic data, and "geodetic" data which includes leveling, trilateration, GPS, InSAR, surface rupture mapping and other data. In other words, "seismic" data contains coseismic timing information and "geodetic" data does not.

The most definitive way to analyze the effects of each data class is to filter out all results that incorporate a class and observe the effect on the slope of the best-fit line. However this presents a problem in that there are only 19 models that do not incorporate any seismic data, and this is not a large enough population to overcome the inherent noise in the data. Instead, we select all models that incorporate at least some seismic data (148 "seismic" models), then all that incorporate at least some geodetic data (60 "geodetic" models), and finally all that incorporate both geodetic and seismic data (41 "joint" models).

The results are shown in Figure 3.9. For a time window of 2 seconds, only the seismic models cause the null hypothesis to be rejected at 95% confidence. By comparison, geodetic and joint inversions show only a slightly positive correlation between early and final slip. This effect is largely due to the size of the seismic model population being 2-3 times larger than the geodetic and joint populations. When 100 random subsamples of 60 models

are taken from the seismic population, the differences between the seismic, geodetic and joint inversions are somewhat diminished. Nevertheless, this result suggests that seismic data biases the inversion toward earlier slip. This is an interesting result, since seismic data contains timing information necessary to properly constrain the early rupture history. This again suggests that we are observing a real effect rather than a product of poor resolution of the early source process.

Figure 3.10 shows the probability of acceptance of the null hypothesis as a function of time window length for the three populations (with the seismic population subsampled 100 times and log-averaged). For time windows between 3 and 7 seconds the null hypothesis is rejected at 95% confidence for all populations. Figure 3.10 also shows that for time windows shorter than 5 seconds seismic data biases the inversions toward deterministic behavior, while for time windows greater than 5 seconds it is geodetic data that biases the models toward deterministic behavior. A related observation is that the minimum probability of acceptance for seismic inversions occurs with a 3 second time window, whereas for geodetic inversions this minimum occurs with a 6 second time window. This may be a consequence of geodetic inversions' increased sensitivity to the slip centroid rather than the hypocenter, which causes slip to be concentrated further down the fault than slip in seismic inversions, in some cases [Wald et al., 1996] eliminating slip at the hypocenter altogether. If this is the case, geodetic inversions may not sample much of the slip on the fault within a shorter time window.

In addition to the diversity of incorporated source data, the 167 models applied over twenty different inversion methods between them, which may account for some degree of variability. The applied method in 31 of the models could not be determined from the source literature for various reasons, and 28 of the models used in-house methods or uncommon methods. The remaining 108 models used one of nine methods. Three of the methods [Cotton and Campillo, 1995; Ide and Takeo, 1997; Ji et al., 2002] are nonlinear inversion methods which invert both for slip amplitude and timing, while the remaining six [Olson and Apsel, 1982; Hartzell and Heaton, 1983; Satake, 1987; Yoshida and Koketsu, 1990; Sekiguchi et al., 2000; Tanioka and Satake, 2001] are linear methods that assume a rupture velocity and invert only for slip amplitude. For simplicity we separate the data into two groups according to this fundamental methodological difference.

When we account for the different population sizes (35 nonlinear vs. 72 linear inversions) we find no statistically significant difference between linear and nonlinear inversions in terms of bias toward deterministic or cascading behavior. This is apparent in Figure 3.11, which plots each population with a 2-second time window. The best-fit lines have similar slopes and the populations have comparable variances, and both cause the null hypothesis of purely cascading behavior to be rejected at the 95% confidence level. The

similarity between linear and nonlinear inversions appears to be robust over time windows up to 10 seconds, and for time windows from 3 to 7 seconds both populations reject the null hypothesis at the 95% confidence level.

3.7. Variability within individual events

Kinematic slip inversions for an event are non-unique, and often different authors will publish very different slip models for the same earthquake. To control for the effect of this variability on our results, we select from the dataset 10 events (Table 3.1) which are represented by 3 or more models and calculate the mean of the final magnitude and the mean of the magnitude within a 2 second time window for each model. We then fit these means in a linear least-squares sense. Although there are 14 events in the dataset with 3 or more models, only 10 have at least 3 models that yield valid measurements within a 2 second time window (nonzero slip and incomplete rupture process), and these are shown in Figure 3.12. We find that the means behave in a comparable manner to the dataset as a whole. If anything the slope is greater than for the dataset as a whole, although that is mostly due to our decision not to exclude outliers in this analysis.

In the context of this analysis we can see that the degree of variability between models (at least in the one parameter we examine) is not the same for different events. Figure 3.12 and Table 3.1 show, for example, that the final magnitude of the 1999 Hector Mine event is very well constrained by four models at $M_w 7.15 \pm 0.03$, but the initial magnitude in a 2 second time window is $M_w 5.81 \pm 0.88$, and varies from $M_w 4.6$ to $M_w 6.8$. Note that the lowest of these data points is considerably outside the range of the other models, but considering the event by itself it makes no sense to strike the lowest data point as an outlier, as it is not much farther from the mean ($M_w 5.8$) than the highest data point. This is consistent with observations that the Hector Mine earthquake was a temporally complex event, potentially with a slow rupture process and long, spatially variable rise times [Kaverina et al., 2002]. Compared to the Hector Mine event, which has well-constrained final magnitude but poorly constrained initial magnitude, the 1994 Northridge and 2000 Tottori events have comparable standard deviations for both measures. The Northridge final magnitude is $M_w 6.70 \pm 0.08$, while the initial magnitude is $M_w 6.02 \pm 0.06$. Similarly, the Tottori final magnitude is $M_w 6.75 \pm 0.06$ and the initial magnitude is $M_w 5.67 \pm 0.09$. For these events the early slip history appears to be more or less consistent between models, though care must be taken in drawing any detailed conclusions from this as we are reducing a 2D heterogeneous slip model to a single number, and much information is lost in this representation.

The observation of strong scaling in this reduced dataset demonstrates that the scaling is robust when using well-studied events. This is true despite the sometimes significant variation among models for the same event [Beresnev, 2003]. It would be impossible to make any reliable observations of scaling if only one model were used for each event, but by using multiple models we are able to overcome this variation to observe robust scaling between initial and final magnitude.

3.8. Conclusions

We began our investigation by including all possible data from the SRCMOD database and 15 additional events, under the premise that the large number of events would overcome the inherent variability in the kinematic inversions. We observe scaling of early slip and magnitude with the final magnitude of these events. To determine whether the scaling is a spurious result of the poor resolution of slip in the early rupture we filtered out those data which are most poorly constrained: models for which the rise time is longer than the time window, and those which had not ruptured more than one grid point within the initial time window. After filtering the data the scaling remains robust, and in fact is more prominent, indicating that poor resolution of early slip is not the cause of the observed scaling. This scaling is also robust for several subsets of the data: inversions derived from seismic data or geodetic data, and joint inversions; nonlinear and linear inversion methods; and well-studied events with multiple models.

If we interpret these findings at face value, we must allow for the possibility that magnitude is at least in part influenced by processes in the early part of the rupture process. Although this is suggestive of at least partially deterministic rupture behavior, it is not inconsistent with individual earthquakes being cascading ruptures. While the individual earthquake may be strongly controlled by the distribution of stress down-fault, it is possible that the conditions encountered by the early rupture affect its ability to propagate past unfavorable conditions down-fault. In this way, while it may be impossible to determine the course any individual rupture will take, the aggregate behavior of many earthquakes can still exhibit magnitude-dependent scaling reflecting the predisposition of ruptures that start out energetically to propagate farther. This interpretation is consistent with the observation that hypocenters tend to be located near high-slip patches for events in the SRCMOD dataset [Mai et al., 2005], as well as with a number of observational studies [Beroza and Ellsworth, 1996; Olson and Allen, 2005; Lewis and Ben-Zion, 2008] that suggest some correlation between the early seismic arrivals of earthquakes and their magnitudes. It is possible that the observed scaling is due to some fundamental assumptions associated with kinematic inversions, which lead to larger slips near the hypocenter for

larger earthquakes. The likelihood of this case is low, as we continue to observe scaling after filtering out those models for which the grid size is large relative to the time window. Assuming that the observed scaling is a real effect of the physics of rupture, it is impossible to characterize, within the framework of this study, the process or processes by which the final event magnitude is influenced by early slip evolution, or even whether the effect is truly deterministic for each individual rupture or only observable in the aggregate behavior of many earthquakes. Nevertheless, the implications of the observed scaling are significant to the wider debate about the nature of earthquake rupture and may have consequences for other questions or applications in seismology.

3.9. Data and Resources

152 of the slip models used in this study are available in the SRCMOD database, available at <http://www.seismo.ethz.ch/srcmod> (last accessed July 2007). Ji's slip inversion for the 2004 Parkfield Earthquake is additionally available at http://www.tectonics.caltech.edu/slip_history/2004_ca/parkfield2.html.

3.10. Acknowledgements

The authors thank Matt Pritchard for providing 15 of the slip models used in this study; Martin Mai for use of the SRCMOD database, and for insightful comments in the early stages of this study; and two anonymous reviewers for reviews that improved this paper significantly. This paper is Berkeley Seismological Laboratory Contribution #09-18.

3.11. Tables and Figures

Initial and Final Magnitudes of 10 Well-Studied Events			
Event	Number of Models	Initial Magnitude	Final Magnitude
1979 Imperial Valley	4	5.24±0.70	6.54±0.07
1989 Loma Prieta	5	6.17±0.20	6.90±0.04
1992 Landers	5	6.18±0.39	7.20±0.04
1994 Northridge	6	6.02±0.06	6.70±0.08
1995 Kobe	7	5.80±0.11	6.88±0.06
1999 Chi-Chi	5	6.12±0.48	7.67±0.04
1999 Hector Mine	4	5.81±0.88	7.15±0.03
1999 Izmit	6	6.17±0.48	7.46±0.07
2000 Tottori	3	5.67±0.09	6.75±0.06
2004 Parkfield	3	5.50±0.09	6.03±0.03

Table 3.1: 10 earthquakes with 3 or more models yielding usable slips within a 2 second time window. 4 additional events in the dataset have 3 or more models, but within a 2-second time window either the event had ruptured to completion, or no slip had yet been observed on the fault plane.

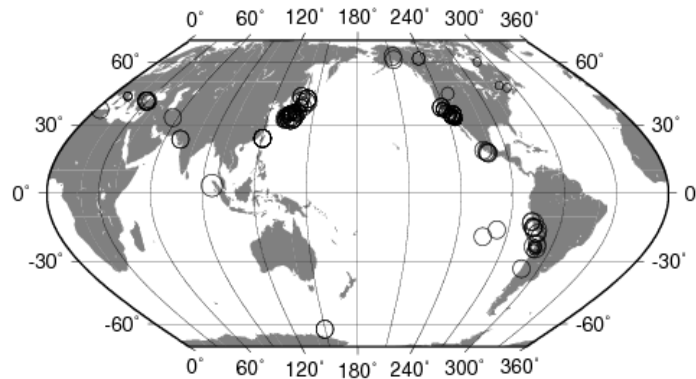


Figure 3.1: Map of the 95 events used in this study. Most of the events occur either in California, Japan, or the west coast of South America.

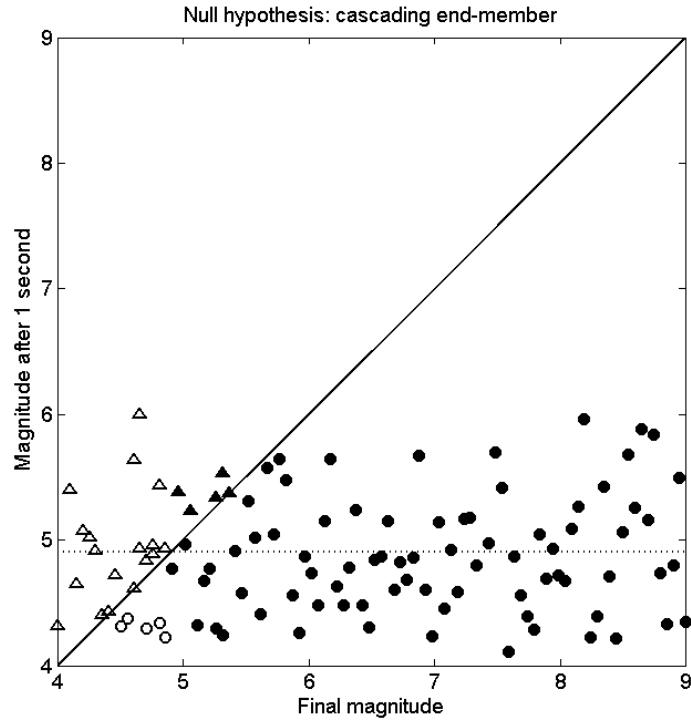


Figure 3.2: The null hypothesis, that earthquakes are cascading ruptures, requires that all earthquakes have approximately magnitude 5 after one second of rupture duration. The dotted line at M_w 4.9 represents the theoretical magnitude after 1 second, and the points are randomly generated and normally distributed about the line. Points above the diagonal line (triangles) are physically impossible. Points with final magnitude less than the reference magnitude (hollow symbols) are culled from the data. The best-fit slope of the remaining points (solid circles) is near zero.

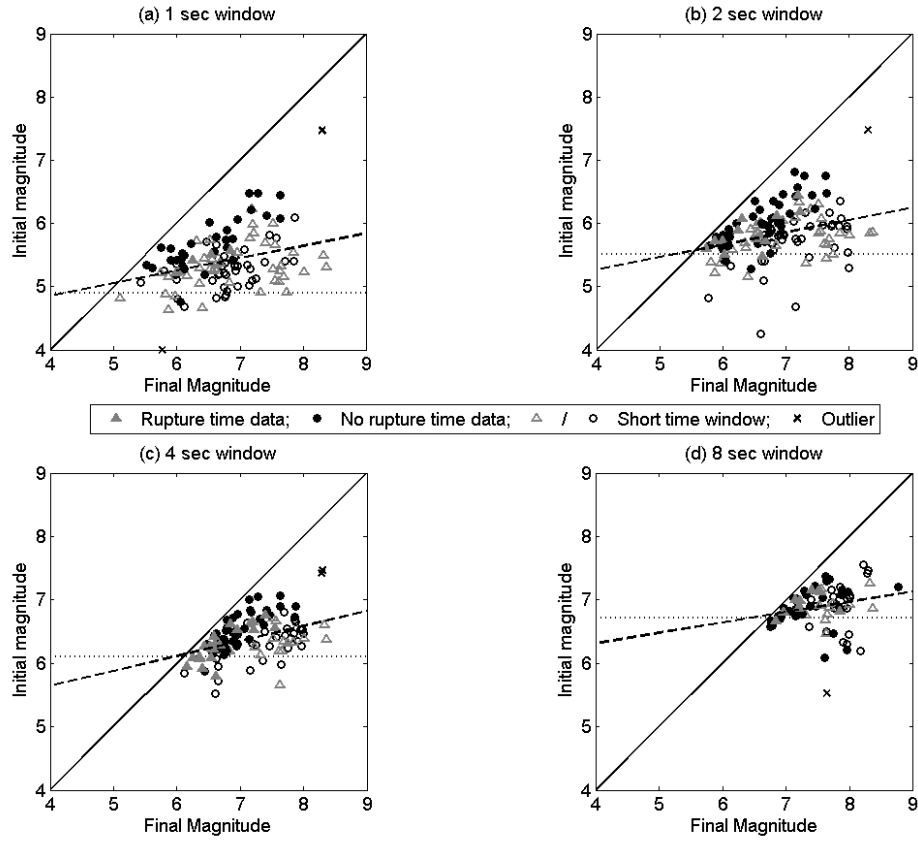


Figure 3.3: Initial vs. final magnitude for time windows of 1, 2, 4 and 8 seconds. Circles are inversions with no point-wise rupture time information, triangles are inversions with rupture time information and hollow symbols are events for which either the time window was shorter than the rise time or only one grid point had ruptured. The dashed line represents the least-squares best fit to the data, and the dotted line is the estimated magnitude for the null hypothesis. Plots (b) and (c) do not show two outliers (crosses) with initial $M \approx 3$. No outliers are included in the fits. For all time windows the null hypothesis can be rejected at the 95% confidence level.

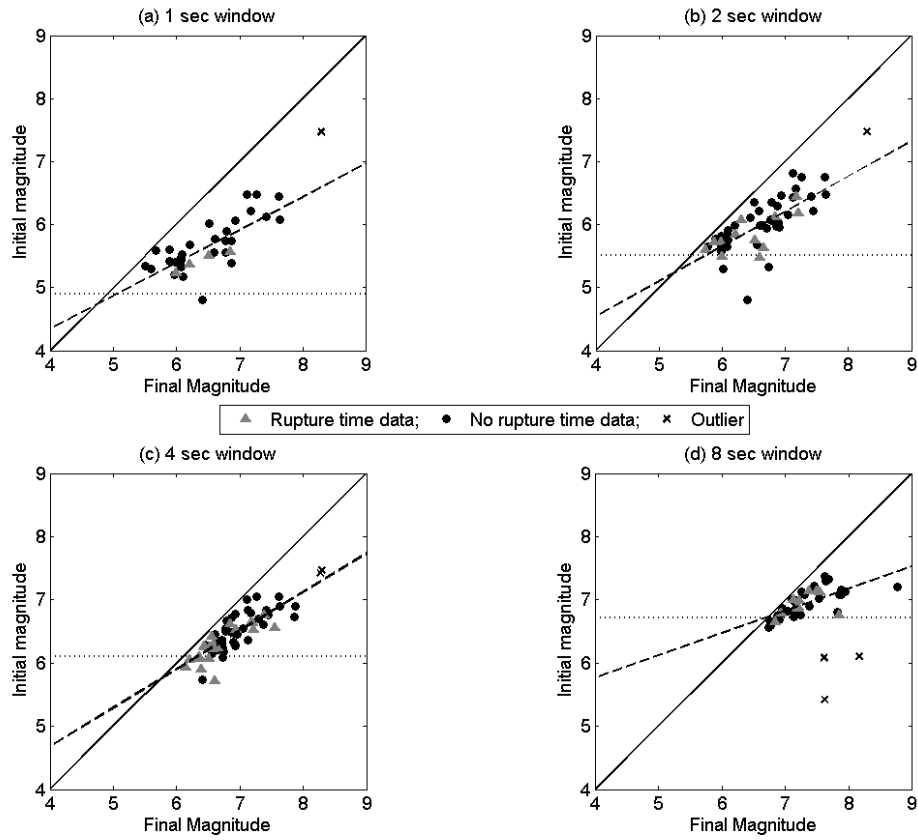


Figure 3.4: Initial vs. final magnitude for time windows of 1, 2, 4 and 8 seconds, after removing points for which the average rise time is less than the time window, or for which only one grid element had begun to rupture within the time window. Magnitudes have been recalculated disregarding any slip less than 10% of peak slip for a given event. The null hypothesis can be rejected at greater than 95% confidence for all time windows.

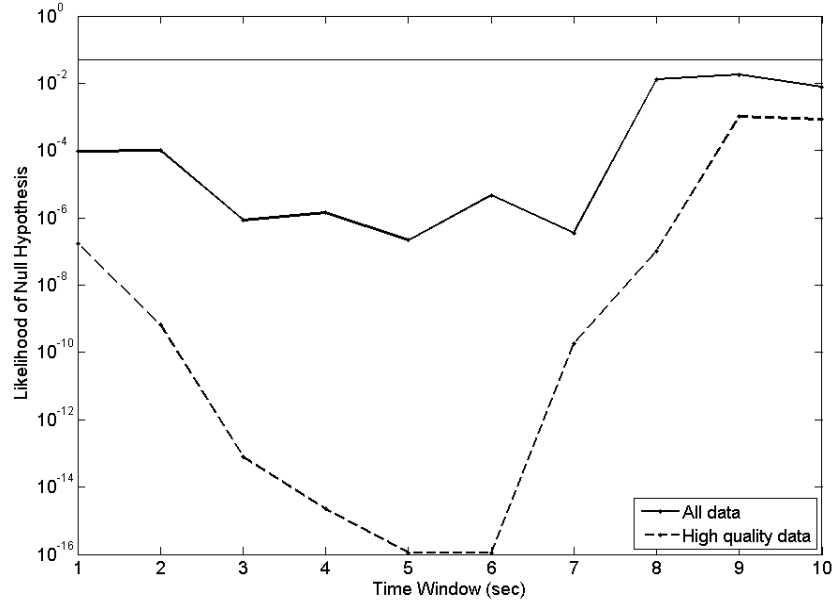


Figure 3.5: Likelihood of acceptance of the null (cascading) hypothesis as a function of time window length. The solid trace is the likelihood using all data (as in Figure 3.3), and the dashed trace is the likelihood using only the high quality data as described in Figure 3.4. Lower likelihood in the dashed trace indicates a more robust scaling of initial vs. final magnitude. The horizontal line indicates the 95% confidence level, and all points below this line reject the hypothesis with 95% confidence.

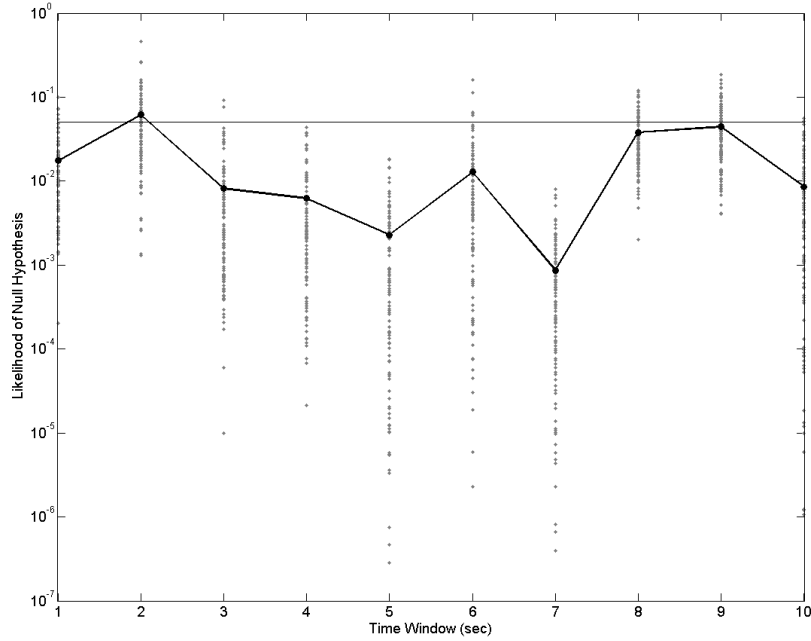


Figure 3.6: Analyses of 100 random subsets of 3 events in each $\frac{1}{4}$ -magnitude bin from magnitude 4 to 8. Gray dots are individual likelihoods of rejecting the null hypothesis, and black circles are the log-average likelihood for all 100 simulations. The null hypothesis can be rejected at 95% confidence for 9 out of 10 time windows

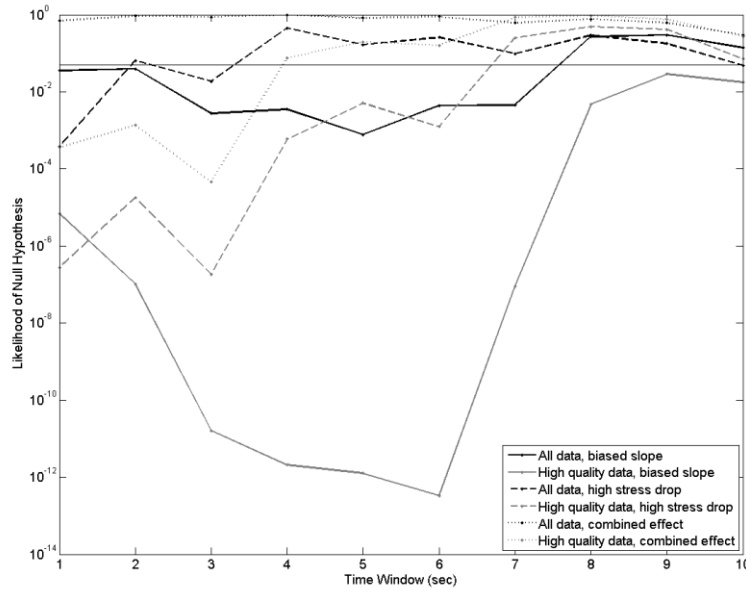


Figure 3.7: Probability of accepting the null (cascading) hypothesis for all time windows using all available data (black lines) and only high quality data (gray lines) as described in Figure 3.4. Solid lines represent the hypothesis test using a null hypothesis slope of 0.105. Dashed lines represent the dataset assuming $\alpha = 10^{-4}$, and a null hypothesis slope of zero. Dotted lines represent the combined effect of choosing $\alpha = 10^{-4}$ and a null hypothesis slope of 0.221. The horizontal line indicates the 95% confidence level, and all points below this line reject the hypothesis with 95% confidence.

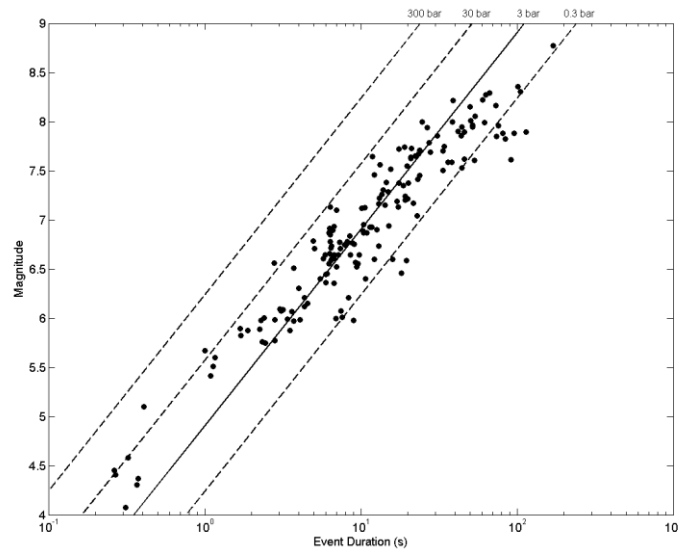


Figure 3.8: Event magnitude plotted against event duration as determined from kinematic slip inversions. Diagonal lines represent theoretical magnitude-duration relations assuming stress drops of 300, 30, 3 and 0.3 bar. The solid line (3 bar) is the closest order of magnitude to the data.

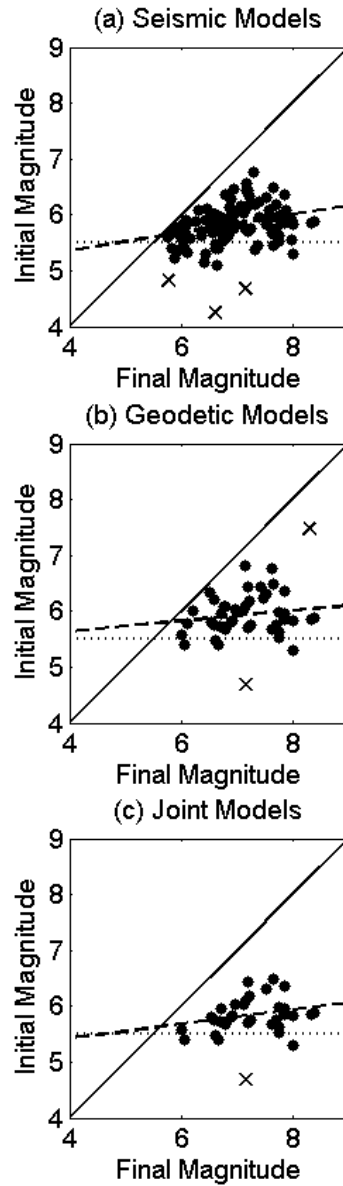


Figure 3.9: Magnitude within a 2 second time window vs. final magnitude for (a) inversions incorporating seismic (strong ground motion or teleseismic) data, (b) inversions incorporating geodetic (trilateration, leveling, GPS, InSAR, surface rupture mapping or other) data, and (c) inversions incorporating both seismic and geodetic data. Crosses are outliers which are not included in the regression.

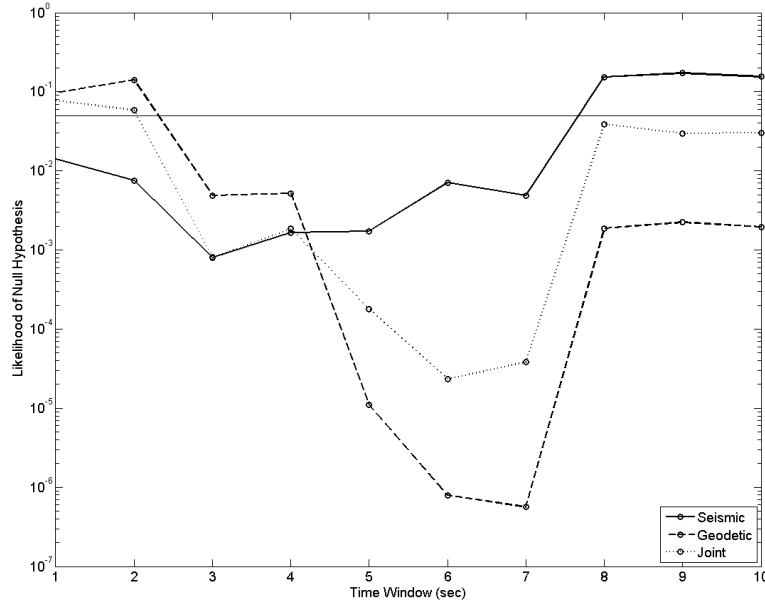


Figure 3.10: Probability of accepting the null hypothesis of cascading behavior for various time windows using models incorporating seismic data, geodetic data, and both. Horizontal rule represents 95% confidence level, below which the null hypothesis must be rejected. For short time windows (4 seconds or less) seismic data biases the inversion toward deterministic behavior. For time windows 5 seconds or longer geodetic data favors deterministic behavior. Seismic probability is the log-average of 100 subsamples of 60 seismic inversions to make the population size comparable to the geodetic and joint populations.

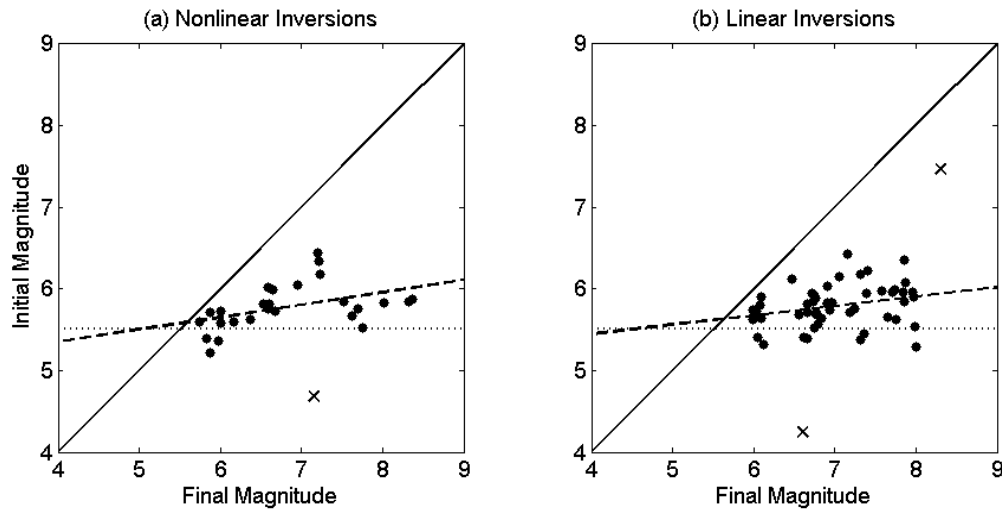


Figure 3.11: Initial vs. final magnitude within a 2 second time window for (a) nonlinear inversion methods [Cotton and Campillo, 1995; Ide and Takeo, 1997; Ji et al., 2002] and (b) linear inversion methods [Olson and Apsel, 1982; Hartzell and Heaton, 1983; Satake, 1987; Yoshida and Koketsu, 1990; Sekiguchi et al., 2000; Tanioka and Satake, 2001]. Crosses are outliers which are not included in the regression. The best-fit relations are statistically indistinguishable from one another.

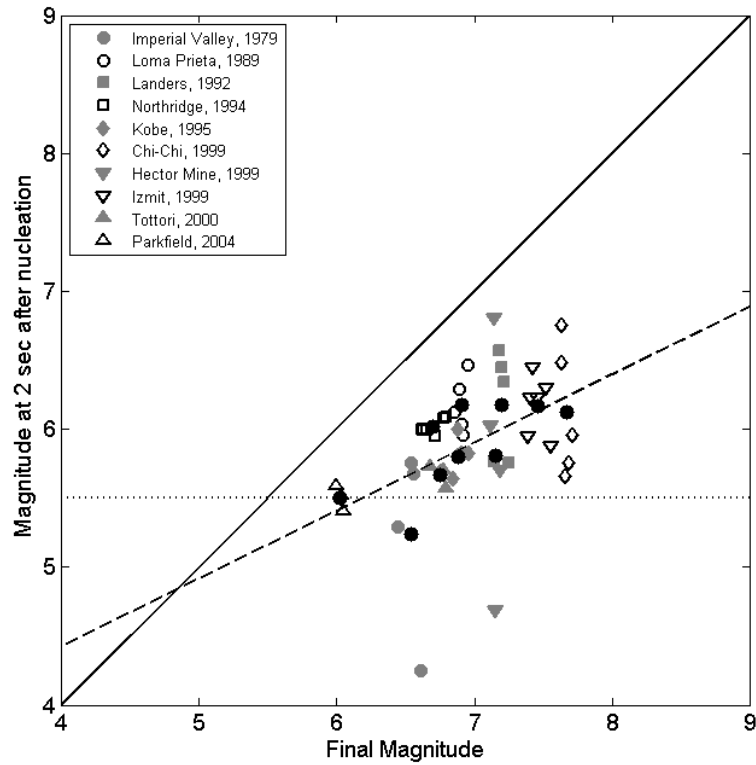


Figure 3.12: Initial vs. final magnitude within a 2 second time window for selected events in the dataset which are represented by more than three models. Black dots show the mean of each event's models, and the best-fit line is based on these means.

4. Modeling the Effect of Early Rupture on Earthquake Magnitude

Gilead Wurman, Richard M. Allen, and David D. Oglesby

4.1. Abstract

Recent observational studies indicate scaling of some properties of early P-wave arrivals with the final magnitude of earthquakes. These studies suggest the possibility that the early rupture history of a given earthquake may influence the behavior of the rupture at later times and farther away on the fault. We test a hypothetical physical mechanism by which such influence might be mediated on the fault plane, by simulating this mechanism using 2D and 3D dynamic fault models and determining whether the hypothesized behavior occurs under realistic fault conditions.

We hypothesize that the intensity of the early rupture, as approximated by the amplitude of shear stress near the focus, imparts more or less energy to the rupture front, and that this intensity either enables the rupture to overcome unfavorable barriers down-fault, or inhibits it from so doing. We use a Genetic Algorithm to search over 11 parameters for the fault properties that are conducive to the transmission of rupture energy down the fault. We find that the hypothesized behavior emerges after only a few search iterations, and examine the fitness function over 9 of the 11 parameters to evaluate whether this behavior is physically realistic.

We find that the fitness function is well-behaved and well-constrained over the majority of the search parameters, and that most of the optimal parameter values agree well with prior studies which address those parameters. We find however that in order for the hypothesized behavior to occur the fault must exhibit significantly stronger velocity-weakening behavior than has been previously observed in the lab. Thermal pressurization or flash heating may cause faults in the field to exhibit velocity-weakening more similar to the values we find, but further modeling incorporating these effects is required to definitively determine whether the hypothesized mechanism is physically realistic. Additionally, several of the parameters appear to need very specific values to support the hypothesized behavior, and this tuning is unlikely in nature. This may be due to the manner in which the search algorithm converges, and further modeling is required to determine whether this is a real effect.

4.2. Introduction

Recent observations of the spectral character of P-waves [*Olson and Allen, 2005; Lockman and Allen, 2005; Lockman and Allen, 2007; Wurman et al., 2007; Lewis and Ben-Zion, 2008*] suggest that there is some information in the early seismic arrivals that may be correlated to the final magnitude of the event, and that in many cases this information is available before the rupture has completed. *Wurman et al. [in review]* find a statistical correlation between early moment release in kinematic slip inversions and their final size. These data pose the question of whether the early rupture

history of an earthquake can influence the development of rupture down-fault, and to what degree that influence extends.

A large number of prior observational studies are at odds over whether small earthquakes are distinguishable from large ones. *Wyss and Brune* [1967] used teleseismic data to distinguish 7 progressively larger subevents making up the 1964 Alaska earthquake, suggesting a cascading rupture. *Abercrombie and Mori* [1994] similarly found that the 1992, M_w 7.3, Landers earthquake was preceded by two subevents, first an M_w 4.4 followed by an M_w 5.6 subevent. *Kilb and Gomberg* [1999] compared the record of an initial subevent of the 1994 Northridge, CA earthquake to records of nearby small earthquakes and found them to be similar, suggesting a lack of scale dependence. However *Umeda* [1990, 1992] found evidence of a region near the earthquake focus which he termed the “bright spot.” He identified early seismic phases in broadband waveforms which he correlated with the formation of the bright spot. His observations indicated that these early phases do scale with the size of the event. *Mori and Kanamori* [1996] concluded that there were no magnitude-dependent precursory phenomena in the 1995 Ridgecrest, CA earthquake sequence. *Ellsworth and Beroza* [1998] observed nucleation phases for the same Ridgecrest events that did scale with magnitude, although they found that these nucleation phases did not argue conclusively for either self-similar or scale-dependent rupture. Similarly, while *Abercrombie and Leary* [1993] used borehole data to examine the source dimensions of earthquakes between M -2 to 8 and found self-similar behavior, *Iio* [1992, 1995] showed that there is an initial slow slip phase that preceded 69 microearthquakes in Japan, and that the duration of this phase is scale dependent. *Beroza and Ellsworth* [1996] found that many earthquakes exhibit an initial phase of low moment rate compared to the rest of the earthquake, which they termed the seismic nucleation phase. They observed this phase for 48 earthquakes of M 1.1 to 8.1 and also found that the size and duration of this phase scale with the event magnitude. *Sato and Mori* [2006] found by applying the model of *Sato and Kanamori* [1999] that the initiation crack size does not significantly vary for events of M 3.5 to 7.9, but *Nakatani et al.* [2000] found that the amplitude of early P-waves for events of M 0.3 to 2.1 does scale with the final magnitude.

Olson and Allen [2005] hypothesized that the influence of early rupture on the size of an earthquake is mediated by the intensity of the early rupture. A higher stress concentration at the focus of the earthquake generates a stronger early rupture phase, and that phase imparts sufficient energy to the rupture to overcome barriers on the fault surface and produce a large earthquake. Conversely, an area of lower stress around the focus generates a comparatively weak early rupture, and the rupture may not gain enough energy to overcome the same barriers, and stops before becoming a large earthquake. Henceforth we use the term “nucleation” to describe the

early dynamic (radiative) rupture history of the event, rather than a long-term (quasistatic) aseismic nucleation. We test the feasibility of this hypothetical model by simulating dynamic ruptures on faults with heterogeneous initial shear stress, varying the stress near the focus while holding the stress constant elsewhere on the fault. We use a Genetic Algorithm to search over a large parameter space for conditions under which these simulations exhibit a correlation between the initial near-focus stress and the final size of the rupture, and test whether these conditions are realistic thus supporting the hypothesis, or they are unrealistic, thereby refuting it.

4.3. Method

We use a multi-domain spectral boundary integral code [Dunham, 2005; Noda et al., 2009] to simulate dynamic fault ruptures in 2D and 3D with heterogeneous initial shear stress conditions (following the work of Ripperger et al. [2007]) on faults with rate-and-state friction laws. The models have some basic properties in common, shown in Table 4.1 (elastic moduli, grid and time spacing, and clamping stress), and differ in 11 parameters (Table 4.2): four controlling the stress distribution (correlation length a_c , Hurst exponent H , mean shear stress τ_μ , and standard deviation of shear stress τ_σ); five rate-and-state parameters (reference friction f_0 , direct-effect coefficient a , evolution-effect coefficient b , characteristic slip-weakening distance L , and reference velocity V_0); and two controlling the size and stress of the overstressed area (also referred to as the “plug”, see Figure 4.4) that initiates rupture (R_{plug} , τ_{plug}). The rate-and-state parameters are used in the friction law [Ruina, 1983]:

$$f = f_0 + a \ln\left(\frac{V}{V_0}\right) + b \ln\left(\frac{V_0 \theta}{L}\right) \quad (4.1)$$

where V is the slip velocity and θ is the state variable, which evolves according to:

$$\frac{d\theta}{dt} = 1 - \frac{\theta V}{L}. \quad (4.2)$$

The large number of parameters to search over makes a direct search of the entire parametric space impossible, so we employ a Genetic Algorithm (GA) scheme [Stoffa and Sen, 1991; Sambridge and Drijkoningen, 1992] to search the parameter space for models that exhibit a dependence of the final rupture size on the early rupture history. We refer to this as “favored” behavior, insofar as we are first attempting to simulate the hypothesized

behavior without regard to realism. The test of this hypothesis occurs when we examine the parametric values found by the GA search to evaluate whether this hypothesized behavior can occur under realistic conditions.

The GA method is designed to emulate the behavior of living populations. An initial population of models is generated by randomly picking parametric values within a search range, and each of the models is tested and evaluated for fitness. When all the individuals of the initial population have been evaluated, the next generation of models is generated by blending the parametric values (the “genes”) of two of the initial models at a time. The higher the fitness of a given model, the more likely it is to be allowed to “breed” into the next generation of models. The fitness function can be simple or quite complex, and in this case the evaluation itself is very computationally expensive. Because GA requires a large population of models to test, it is not as efficient at finding the optimal solution as a Simulated Annealing method [Liu et al., 1995; Beaty et al., 2002]. As we explain below, we are less interested in the single optimal solution than we are in the behavior of the fitness function over the range of possible values for each parameter. Simulated Annealing only evaluates one model at a time, and would not sample the model space as well as GA, so we choose to employ the latter method.

The GA search process is illustrated in Figure 4.1 through Figure 4.3. Each generation of models has 30 individual models (“individuals”, Figure 4.1), and each individual is used to generate 10 stochastic stress distributions (“realizations”, Figure 4.2) using the method of *Ripperger* et al. [2007] and the parameters a_c , H , τ_μ and τ_σ (Table 4.2). This yields 300 different fault models with 300 different shear stress distributions. We then find the point of greatest shear stress on each fault and set that to be the point of nucleation by overstressing it with a cosine function of radius R_{plug} and magnitude τ_{plug} and allowing the model to rupture spontaneously. We approximate varying the “intensity” of nucleation by taking the stress in the nucleation region, defined as the area surrounding the nucleation site, within the contour defined by the cutoff stress:

$$\tau_{cutoff} = C(\tau_{max} - \tau_{min}) + \tau_{min} \quad (4.3)$$

where C is a constant equal to 0.5 for the 2D case and 0.65 for the 3D case. We scale the nucleation region according to the relation:

$$\tau_{scaled} = \tau_{cutoff} + K(\tau_{orig} - \tau_{cutoff}) \quad (4.4)$$

with K varying in 10 equal increments from 10% to 100% (Figure 4.4). Each of these increments is referred to henceforth as a “step.” Note that for

extremely small values of a_c and large values of R_{plug} , it is possible that the overstressed region (the “plug”) extends beyond the nucleation region bounded by τ_{cutoff} . There is no constraint on this condition within the search algorithm, but in practice we find that models with such short correlation lengths do not perform well, and the GA search selects against such models. In all the models that exhibit favorable behavior, the overstressed plug is completely contained within the nucleation region.

We choose to use GA with elitism, meaning the 2 best individuals from the previous generation are copied exactly to the next generation, to avoid losing the best solution in the next generation. The individuals are bred using roulette selection, where the probability of an individual breeding is proportional to its share of the total fitness of the population. When two individuals (A and B) breed they produce two children (a and b) with complementary combinations of the parents’ genes (Figure 4.3). It is typical in GA methods to assign a likelihood (termed the crossover probability) that parent A’s genes will end up in child b’s genome and vice-versa. This probability ranges from about 60% to 100% [Stoffa and Sen, 1991; Sambridge and Drijkoningen, 1992], and we find that our algorithm converges fastest for simple problems with a crossover probability of around 85%. The results of the search are not affected by this probability, only the rate at which the search converges. Most GA schemes use a binary genome (i.e., a given gene is either “on” or “off”), but we choose to use a numeric genome where genes can have any value within a range. Therefore we have to consider two types of heredity: a direct inheritance of one or the other of the parents’ genes, and a blended inheritance of some linear combination of both parents’ genes. Given a crossover, we allow an even chance that a gene will be directly inherited or blended. Thus, there is a 15% probability that parent A will pass a gene to child a, a 42.5% chance that parent B will pass the gene to child a, and a 42.5% chance that child a’s gene will be a linear combination of the genes from both parents. Child b’s genes are the complement of child a’s, so if child a receives a gene from parent B, child b will take the same gene from parent A. If child a receives a gene that is 40% from parent A and 60% from parent B, then child b will have 60% from parent A and 40% from parent B, and so on. There is also a 4.55% chance that a given gene will mutate to a random value within the search space. This probability yields an average of one mutation in every other individual (1 over 11 genes times 2 individuals), and prevents the GA search from falling into a local minimum in the fitness function.

The fitness function itself tests 8 properties of each realization, summarized in Table 4.3. If the result of the test is positive, the value of each test is added to an individual’s score. If the test result is negative, the score is not changed. Positive values indicate favored behaviors, and negative values indicate disfavored behaviors. For example, if the rupture propagates beyond the nucleation region (Figure 4.4) this is considered

favorable behavior (though not extremely so), while if the rupture size grows by more than 80% in any stress step this is considered strongly favorable behavior, and if such growth occurs between two or more steps, it is considered very favorable behavior. On the other hand, any rupture propagating at greater than the P-wave speed is considered unphysical and therefore strongly unfavorable. If the rupture never becomes self-sustaining (i.e., never propagates beyond the overstressed plug) this is also unfavorable behavior. If the rupture is still ongoing at the end of the simulation time, or if the entire model has ruptured by the end of the simulation, we consider this to be a runaway rupture, which is a somewhat unfavorable behavior (the entire model should never rupture within the simulation time unless there is super-P rupture velocity, but the check is included as a safeguard). If the rupture jumps forward and initiates a secondary rupture ahead of the primary rupture front, this is considered slightly unfavorable behavior. This is not because the behavior is unphysical, but because it adds complexity to the evaluation of the model by a computer. There is an additional penalty of -30 to models that failed to run altogether due to the parameter combination causing the simulation to become numerically unstable. Each realization can get a score between -30 and +22, though typical scores range from -15 to +15. The total score of each individual is the sum of the scores of its 10 realizations, and typically ranges from -150 to +150. This score is normalized by a factor of 150, and the individual's fitness is 10 to the power of this value. This yields positive values of fitness ranging over approximately 2 orders of magnitude. Since each individual starts with a score of 0 (i.e., a fitness of $10^0 = 1$), any individual with favorable traits will have a fitness greater than 1, and any unfavorable individual will have a fitness less than 1.

This method introduces a bias to the model output, in that we are explicitly selecting for behavior which shows a correlation between initial shear stress and final size of the event. For this reason, the existence of models that exhibit this behavior is necessary but not sufficient to support the hypothesis. To further support the hypothesis we evaluate the resulting parameter values themselves in three ways:

- Regularity: is the fitness of the models a smooth, well-behaved function of a given parameter, or does the fitness vary irregularly as the parameter varies?
- Constraint: do high-fitness models exhibit a wide range of parametric values, or are they constrained to only a narrow range of values for the given parameter?
- Correctness: do high-fitness models exhibit the same or similar parametric values to those found by preceding field and lab studies, or are they very different?

In general, the more regular the fitness function, the better constrained, and the more the parametric values agree with previous studies, the more well-

supported the hypothesis becomes. In particular, the values of the rate-and-state parameters are extensively studied and provide a good reference point for comparison. That is, if we see models that exhibit the favored behavior only under unrealistic frictional conditions, this is a strong argument against the hypothesis.

4.4. Modeling in 2D

We ran the GA search over 20 generations (of 30 models each) in a 2D space, on a line fault of 400 km length with a grid spacing of 100 m. Although a line fault behaves differently from a planar fault in a full 3D space, this preliminary step is important for several reasons. First, it verifies that the search algorithm converges to the behavior of interest in a reasonable amount of time, and allows us to adjust the fitness scoring algorithm relatively quickly. The 2D model search completes one generation within 20 hours or less, whereas the 3D model search takes as many as 11 days to complete a single generation. By starting with the 2D search we are able to set up and debug the scoring algorithm over the course of a few days rather than a few months.

Second, the 2D models can be larger and run for longer simulation times than the 3D models, while still remaining computationally tractable. Because the MDSBI code requires periodic boundary conditions, the model size must be significantly larger than the size of the desired ruptures. For a 400 km model, the P-wave propagates around the model in 33 seconds, allowing for a 30 second simulation. By contrast, the 3D models can only be 60 km by 60 km, with only about a 6 second simulation. The use of 2D gives the rupture greater opportunity to propagate beyond the influence of the artificial nucleation, and in general allows for a greater variation in scale than do the smaller 3D models.

Finally, the optimum set of parameters from the 2D search can be used as a starting point for the more computationally expensive 3D modeling. This bootstrapping approach potentially saves several generations and many weeks of search time and allows the 3D search to converge almost immediately.

Within about 16 search generations the 2D models begin to exhibit a dependence of final magnitude with the stress level of the nucleation region. This behavior is widespread after 20 generations. When considering all 20 generations plus the zeroth generation (630 individuals in total), 23.5% have fitnesses greater than $10^{0.5}$, indicating significantly favorable behavior. Figure 4.5 shows a typical example of rupture size varying with stress level through 4 out of the 10 steps of a particular stress realization. The actual parameter values for this individual are listed in Table 4.4. The black traces show the rupture time of each point on the fault with units on the left x-axis. The blue traces show the initial shear stress on the fault, with units on the

right x-axis. When the stress is below about 30% of maximum (step 3), only the nucleation asperity ruptures. This is about 4 km of the length of the fault and can be seen in the upper left of Figure 4.5. When the nucleation stress is scaled to 50% of maximum (step 5, upper right of Figure 4.5) an additional asperity ruptures to the right of the nucleation. When the nucleation is scaled further to 80% (lower left of Figure 4.5) an additional asperity ruptures, and finally at 100% of the original shear stress, a fourth asperity ruptures (lower right of Figure 4.5).

It is clear from Figure 4.5 that at least in one example the final magnitude of an earthquake is strongly dependent on the amount of stress encountered by the early rupture, even when stress on the rest of the fault is held constant. However, this example shows a single realization of a single individual in a single search generation. The question remains whether this is a widespread behavior even among models with the same parameter values as the model in Figure 4.5. We can assess this by comparing the final magnitude of each event to the magnitude of the same event after some fraction of the rupture duration. Figure 4.6 shows the final magnitude of 100 realizations of stress using the parameter values of the individual represented in Figure 4.5 and Table 4.4. The final magnitude is plotted against the initial magnitude after 0.4 seconds of rupture. Each realization is represented by 10 points, one for each stress step.

If the models in Figure 4.6 were insensitive to the stress level in the nucleation region, these points would have a slope around zero. Note that the blue data points represent models that release 97% or more of their final magnitude within the first 0.4 seconds. Because the initial magnitude in these models is effectively saturated (initial magnitude can never exceed final magnitude) we exclude these data from the least-squares regression shown in Figure 4.6. We calculate this slope using only the black data points to avoid artificially increasing the slope of the best-fit line. As is apparent from Figure 4.6, we determine that the slope of the points is greater than zero with 95% confidence using a one-sided t-test. This indicates that the final magnitudes of the individual models are responding nonlinearly to small perturbations in the initial magnitude. This behavior is in agreement with results from the analysis of kinematic slip distributions [*Wurman et al.*, in review], i.e. a large difference in final magnitude corresponds to an observable but smaller difference in the early moment release. The strongly positive slope of the data indicates that the behavior observed in Figure 4.5 is quite common.

4.5. Modeling in 3D

We take the 20th generation results from the 2D modeling and apply them as the “zeroth” generation for modeling the behavior in 3D. Because the parameter values by the 20th generation are quite homogeneous, we take 9

individuals from this generation with scores less than $10^{0.5}$ and replace them with new individuals with random parameter values from within the search space. This is done to allow for the possibility that the best parameter values in 2D are not the same in 3D, but as a practical matter we observe that these new individuals perform quite poorly, and are selected out within 2 generations. This suggests that the choice of bootstrapping from 2D to 3D is appropriate, and reduces (though it does not eliminate) the possibility of a better solution somewhere in the search space.

As with the 2D modeling, each search generation is composed of 30 individuals with 10 iterations per individual. Due to limits on computational power and the added dimension, the models are reduced to 60 km x 60 km at 100 m grid spacing, rather than 400 km on a side as in the 2D case. Due to the smaller fault size the models can only simulate 6.32 seconds of rupture time, at which point the P-wave begins to intersect the S-wave across the periodic boundary conditions, potentially affecting the rupture in an unrealistic fashion. The cutoff stress coefficient C in Eq. 4.3 is 0.65 rather than 0.5 as in the 2D case. This is necessary because the connectivity between grid points is greater on a 2-dimensional fault plane, so that if the cutoff is set as in the 2D case the contour often encloses half or more of the fault plane. On the other hand, if the cutoff level is set too high our ability to affect the intensity of nucleation is constrained, since we can only vary the stress in this region between its original value and the cutoff value. By setting the cutoff at 0.65, we strike a balance between small, contained nucleation regions and having enough ability to affect the nucleation intensity.

Because the parameter values for the 3D search are bootstrapped from the 2D results, the models exhibit a rupture size that is dependent on the early stress magnitude from the zeroth generation. Out of 180 individuals in five search generations plus the zeroth generation, 42.2% exhibit significantly favorable behavior, defined as fitness greater than $10^{0.5}$. In general the optimal parameter values in the 3D search are unchanged from the 2D search because of this bootstrapping, but the parameters that are less well-constrained in the 2D search do migrate somewhat in the 3D search. Notably the correlation length a_c becomes somewhat shorter in the 3D search, as does the critical slip-weakening length L . This is discussed in more detail in the following section. Figure 4.7 shows an example of a realization that exhibits a dependence of rupture size on the initial stress level. The colors indicate the initial shear stress on the fault, warm colors indicating high stress asperities (up to ~ 50 MPa) and cool colors indicating barriers. The superimposed contours show rupture time in 0.2 second intervals. Similar to Figure 4.5, when the nucleation stress is less than about 20% (upper left of Figure 4.7) only the nucleation asperity ruptures, but as the stress is scaled up to 40% (upper right of Figure 4.7), 70% (lower left), and finally 100% (lower right), a second, third and fourth asperity rupture

as well. The specific fitness and parameter values of this individual are shown in Table 4.4.

4.6. Comparison of Parametric Values

It is likely that a particular parameter value is necessary to a high-fitness model but not sufficient by itself. Therefore the fitness when plotted against a single parameter may have a great range of values for a single value of the parameter. If we take the maximum fitness over a narrow range of values for the given parameter we can identify those values which allow for high-fitness behavior as opposed to those values which preclude it. Figure 4.8 shows the individual value/fitness pairs for the nine parameters controlling friction and stress distribution, with lines representing the maximum fitness at each value for the parameter. Black points and lines show the 2D results, while blue points and lines show the 3D results. Note that we examine the quantity $b - a$ rather than b , as the difference between the state variable and the rate variable describes the degree to which the fault friction is velocity-weakening [Tullis and Weeks, 1986; Marone, 1998; Scholz, 1998; Rice et al., 2001]. As stated in the previous section, the optimal parameter values for the 3D search are very similar to those for the 2D search because they were bootstrapped from the last generation of the 2D search. This is particularly true of the well-constrained parameters τ_{μ} , τ_{σ} , f_0 , V_0 , a and b (Figure 4.8c-h). However there is some migration in the values of the parameters a_c (Figure 4.8a) and L (Figure 4.8i) to shorter length scales. The correlation length a_c shows good fitnesses in the 2D search from 250 km down to about 50 km. The 3D search shows good fitnesses around 50 km but not at longer correlation lengths, and exhibits a new peak around 30 km as well. This is likely due to the physical extent of the fault being much smaller in the 3D search: only 60 km on a side as opposed to 400 km for the 2D search. This makes correlation lengths longer than about 50 km (i.e., longer than the fault itself) unusable as there is only one asperity and one barrier on the entire fault plane. The same trend is observed in the slip-weakening length L , with the 3D search losing fitness at slip-weakening lengths greater than about 18 cm, and gaining a new peak around 7 cm. Contrary to the correlation length, it is not clear why this should be true for the slip-weakening length. There is no obvious physical reason for shorter slip-weakening lengths to be favored on the smaller faults, since all the length scales are much smaller than the size of the fault anyway. The grid spacing is also constant at 100 m for both the 2D and 3D spaces. The physical significance of this result is not clear, and further study may be necessary to shed light on it.

The results outlined in the previous sections are necessary but not sufficient to support the hypothesis. While demonstrating that a relationship between early rupture history and final magnitude is physically possible for

some combination of parameters, the observations do not conclusively confirm the hypothesis. To test the hypothesis in a falsifiable fashion, we ask the three questions outlined in the Method section:

- Whether the fitness surface for the search is well-behaved (regularity)
- Whether those parameter values are very narrowly defined (constraint)
- Whether the observed behavior occurs under realistic parameter values (correctness)

In this analysis we ignore the two model parameters R_{plug} and τ_{plug} as they have no physical significance outside this modeling exercise. The question of regularity can be applied to all nine physical parameters, but the questions of constraint and correctness are best applied to parameters for which we have some knowledge of the range of accepted values. Of the nine parameters, only the mean shear stress and the five rate-and-state parameters (a , b , f_0 , V_0 , and L) can be compared to values found in previous field and lab studies.

4.6.1. Regularity

Because of the high dimensionality of the parameter space, the regularity of the fitness surface in any single dimension is very difficult to evaluate. The fitness surface will only be regular along any parameter if that parameter exerts a very strong influence on the overall fitness of the model, to the point where it dominates the signal from the other dimensions. Thus an instructive way to inspect the regularity is by considering the maximum fitness as illustrated by the solid lines in Figure 4.8.

The stress-related parameters a_c and H (Figure 4.8a and b) show very irregular behavior, with many peaks and valleys in the maximum fitness. The same is true of the slip-weakening distance L (Figure 4.8i), but the six remaining parameters (Figure 4.8c-h) behave reasonably well, with only one major peak in the fitness curve, and one or two smaller peaks that are typically supported by only one individual. While the fitness surface is well-behaved in the majority of the dimensions, the irregularity along the three parameters a_c , H and L requires further examination. In this examination it is instructive to first look at the degree of constraint of these and other parameters.

4.6.2. Constraint

It is immediately apparent that the parameters a_c , H and L are also the least constrained of all the parameters, as they have the broadest range over which high fitness values can be found. In particular H appears to be equally favored across the entire search space, except possibly for values greater than approximately 2.2. In this context it is possible that some of the irregularity observed in these three parameters is due to the fact that not all values of a_c , H and L occur alongside favorable values of the other parameters, though there are too few data points to test this in a statistically

meaningful fashion. If this is the case, however, it is likely this would lead to the observed narrow valleys in the fitness curves (Figure 4.8a, b and i). In contrast, broader valleys might be expected if there were values of these parameters which actually discouraged high-fitness behavior. It is important to note that the three parameters which are least well-constrained in the Genetic Algorithm search are also those which are least well-constrained observationally (as detailed in the next section).

The remaining six parameters (Figure 4.8c-h) are quite well constrained, to the point where the high-fitness peaks for τ_{μ} , f_0 , a and $b - a$ are nearly delta functions. This may be due to the manner in which the Genetic Algorithm converges on the best solution, or it may be due simply to the fact that these parameters have to be very finely tuned to enable the nucleation energy to influence the rupture over long distances. If the latter is the case this would argue against the hypothesis, as it is unlikely that the mean shear stress in the crust is so homogenous, or that the friction parameters a and b are uniform among varied rock types. A directed search around these parameters is required to determine whether this fine tuning is really necessary.

4.6.3. Correctness

Of the nine parameters plotted in Figure 4.8 only the latter six are at all constrained by observational data. The shaded regions in Figure 4.8d-i show the approximate range of values for mean shear stress [*Choy, 1995; Mayeda and Walter, 1996; Ide and Beroza, 2001*] and the five rate-and-state parameters [*Dieterich, 1978; Rice, 1983; Tullis and Weeks, 1986; Scholz, 1988; Blanpied et al., 1991; Dieterich and Kilgore, 1994; Dieterich and Kilgore, 1996*]. The parameter L is rather poorly constrained in nature (Figure 4.8i), varying from the scale of approximately 50 cm in the field [*Ide and Takeo, 1997; Mikumo et al., 2003; Wibberley and Shimamoto, 2005*] to microns in the lab [*Tullis and Weeks, 1986; Ohnaka and Shen, 1999*]. The mean shear stress on real faults (Figure 4.8d) has not been directly observed except for in a mere handful of cases (e.g. the SAFOD borehole, [*Hickman and Zoback, 2004; Tembe et al., 2010*]), but its value can be constrained to be on the order of 3 to 30 MPa from indirect observations of seismic stress drop in earthquakes [*Choy, 1995; Mayeda and Walter, 1996; Ide and Beroza, 2001*]. The values of f_0 and a and $b - a$ (Figure 4.8e, g and h) are principally constrained by laboratory studies [*Dieterich, 1978; Rice, 1983; Tullis and Weeks, 1986; Scholz, 1988; Blanpied et al., 1991; Dieterich and Kilgore, 1994; Dieterich and Kilgore, 1996*] and are not directly observable in the field either. The value of V_0 (Figure 4.8f) is generally taken to be approximately equal to plate motion rates. As there are no direct observations of stress on faults other than at a few disparate points, there is no constraint on the spectral character of stress distribution (a_c and H), or of the magnitude of its variation (τ_{σ}).

Four of the six parameters are optimal at realistic values. The fact that V_0 is about one order of magnitude too large is of some concern, but in fact V_0 is a reference velocity and has limited physical significance, since the values of a and b can be adjusted slightly (less than 0.2 log units) to compensate for this and achieve the same friction with more realistic V_0 . Of greater concern is the significant difference between observed values of $b - a$ and the optimal value in Figure 4.8h. This suggests that in order for the fault to display the observed behavior it must be much more strongly velocity-weakening than is commonly observed in the lab. There has been some recent work on the effects of flash heating of fault materials and thermal pressurization of pore fluids [Wibberley and Shimamoto, 2005; Noda et al., 2009], both of which may increase the degree of velocity-weakening on real faults as opposed to lab studies. However, there is no analytical solution for the effect these mechanisms would have on our results. Absent such a solution, the stronger-than-expected velocity-weakening argues against the hypothesis being realistic, though not definitively so.

4.7. Conclusions

Using a Genetic Algorithm search over 11 parameters we were able to find a set of models which exhibit a correlation between the stress encountered by the early rupture and an event's final magnitude. This behavior exists both in a 2D space on a linear fault, and in 3D space on a planar fault. In general the results occur with the same parameter values in both 2D and 3D space, though there is a slight shortening in the correlation length and the slip-weakening distance for faults in 3D. This may be an artificial product of using smaller faults in 3D. The Genetic Algorithm fitness curve is largely well-behaved and the optimal parameter values are well-constrained. 630 individual models were tested in the 2D case, and 180 models were tested in the 3D case. In each case, a large proportion of the individuals exhibit favorable behavior to some degree (23.5% in the 2D case, and 42.2% in the 3D case). For those parameters which have been addressed in prior studies, the optimal values we find are largely in agreement with these prior studies. The exception is the degree of velocity-weakening in the friction law, expressed in the value of $b - a$ in the rate-and-state friction law of *Ruina* [1983]. We find that much stronger velocity-weakening behavior is needed than is found in lab studies of typical fault rocks. This presently argues against the hypothesis being realistic, but further modeling is needed which incorporates the effects of thermal pressurization and flash heating, which may reduce the necessary velocity-weakening to realistic values. Additional modeling must be done to determine if several of the parameters must be very finely tuned in order to support the hypothesized behavior, as such fine tuning is unlikely in nature.

In general, the correlation of the final magnitude with some parameter of the early rupture history is consistent with various observations that certain properties of early P-waves vary with earthquake magnitude [*Olson and Allen, 2005; Lockman and Allen, 2005; Lockman and Allen, 2007; Wurman et al., 2007; Lewis and Ben-Zion, 2008*], but whether the behavior modeled in this study is a realistic explanation of the seismic observations remains unanswered. Also unaddressed is the specific mechanism by which the intensity of the early rupture phase is modified, since the artificial scaling of near-focus stress is simply a proxy for that intensity. Finally, it is necessary to demonstrate that whatever physical mechanism drives the intensity of nucleation also affects the properties of the P-wave in a manner consistent with seismic observations. However, the behavior hypothesized in these models is an important first step toward positively identifying the physics that generate these seismic observations.

4.8. Acknowledgements

The authors would like to thank Eric Dunham for contributing the MDSBI code and assisting in its use, and for many helpful discussions on rate-and-state friction.

4.9. Tables and Figures

Physical Parameters Common to All Models		
Parameter	Symbol	Value
Rigidity	μ	$3.204 \times 10^{10} \text{ N/m}^2$
S velocity	V_s	3464.1 m/s
Clamping stress	σ_n	120 MPa
Grid spacing	dx	100 m
Time interval	dt	0.008 sec

Table 4.1: Physical parameters of modeled faults and simulation parameters. These do not vary from model to model.

Parameters in Search Space		
Parameter/gene	Symbol	Range
Correlation length	a_c	$[10^0, 10^{3.5}] \text{ m}$
Hurst exponent	H	$[-1, 3]$
Mean shear stress	τ_μ	$[10^6, 10^{8.5}] \text{ Pa}$
Stress deviation ratio	τ_σ/τ_μ	$[10^{-2}, 10^0]$
Baseline friction	f_0	$[0.4, 0.8]$
Rate coefficient	a	$[10^{-3}, 10^{-1}]$
State coefficient ratio	b/a	$[10^0, 10^2]$
Critical distance	L	$[10^{-1.5}, 10^{0.5}] \text{ m}$
Reference velocity	V_0	$[10^{-13}, 10^{-3}] \text{ m/s}$
Plug radius	R_{plug}	$[100, 1500] \text{ m}$
Plug stress ratio	τ_{plug}/τ_σ	$[10^{-1}, 10^1]$

Table 4.2: Genes in the Genetic Algorithm genome, their symbolic representation, and the range of values they are allowed to take in the search process. Ranges indicated as a power of 10 indicate that the parameter is encoded in log form in the search.

Behavioral Tests in Fitness Evaluation	
Test	Value
Rupture propagated out of nucleation zone	+2
Final size 80% greater than previous stress step	+5
Multiple size increases in the same realization	+15
Fault still rupturing at end of simulation	-6
Super-P rupture speed	-15
Entire model is ruptured	-6
Rupture never became self-sustaining	-10
Secondary rupture initiates ahead of rupture front	-3

Table 4.3: Fitness tests to which model output is subjected and value associated with passing each test. Positive values represent desirable behaviors, and negative values are undesirable behaviors.

Parameter values for example individuals		
Parameter	2D (Figure 4.5)	3D (Figure 4.7)
a_c	53.8 m	31.7 m
H	-0.095	0.414
τ_μ	21.5 MPa	21.5 MPa
τ_σ/τ_μ	0.247	0.278
f_0	0.651	0.651
a	0.0054	0.0054
b/a	6.50	6.45
L	15.4 cm	6.94 cm
V_0	0.155 $\mu\text{m/s}$	0.162 $\mu\text{m/s}$
R_{plug}	462 m	642 m
τ_{plug}/τ_σ	4.02	4.45
<i>Fitness</i>	13.6	13.6

Table 4.4: Values of the search parameters and fitness of the two individuals represented in Figure 4.5 (2D case) and Figure 4.7 (3D case). The values for the 2D case are also used to generate Figure 4.6.

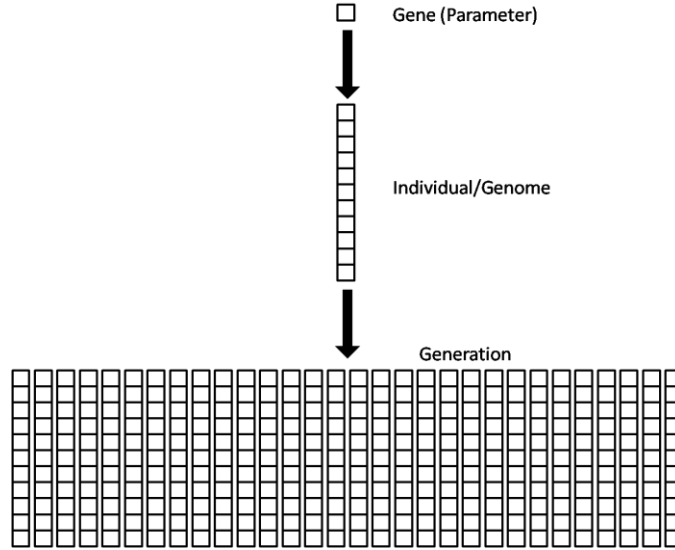


Figure 4.1: Nomenclature used in the Genetic Algorithm search. Each parameter is represented by a single gene, and the 11 different genes together make up the complete genome of a single individual. Each generation of the search is composed of 30 individuals with different genomes. The “zeroth” generation is made up of randomly selected parameter values, and subsequent generations have parameter values derived from the previous generation’s genes.

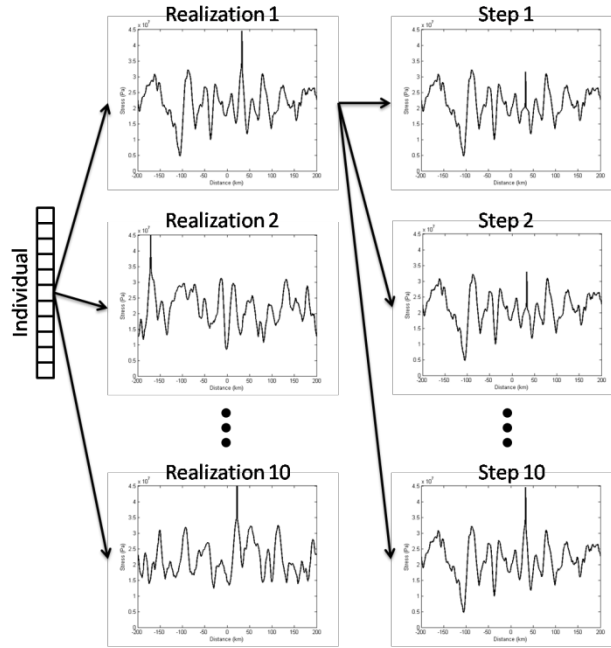


Figure 4.2: Illustration of the derivation of models from individual genomes. Each individual is used to create 10 random realizations of stress using the genes for the 4 stress parameters a_c , H , τ_{pl} and τ_{σ} . A cosine-shaped plug of radius R_{plug} and amplitude τ_{plug} is superimposed on the point of highest stress in each realization. For each realization, this nucleation asperity is then scaled down in 10 equal steps as described in Figure 4.4. All realizations and steps have the same rate-and-state friction law defined by the genes for the 5 friction parameters f_0 , a , b , L and V_0 .

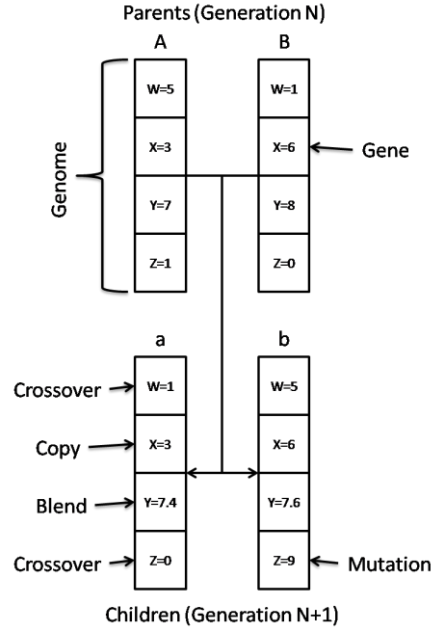


Figure 4.3: Illustration of the ways in which genes get passed from parents to children. Gene X is copied directly from each parent: child a gets the X gene from parent A, and child b gets the X gene from parent B. Gene Y is blended: child a gets 40% of the Y gene from parent A and 60% of the Y gene from parent B, while child b gets 60% of parent A's gene and 40% of parent B's gene. Genes W and Z are copied directly from the respective parents with no alteration, but a random mutation has occurred in the Z gene of child b, and its value is randomized.

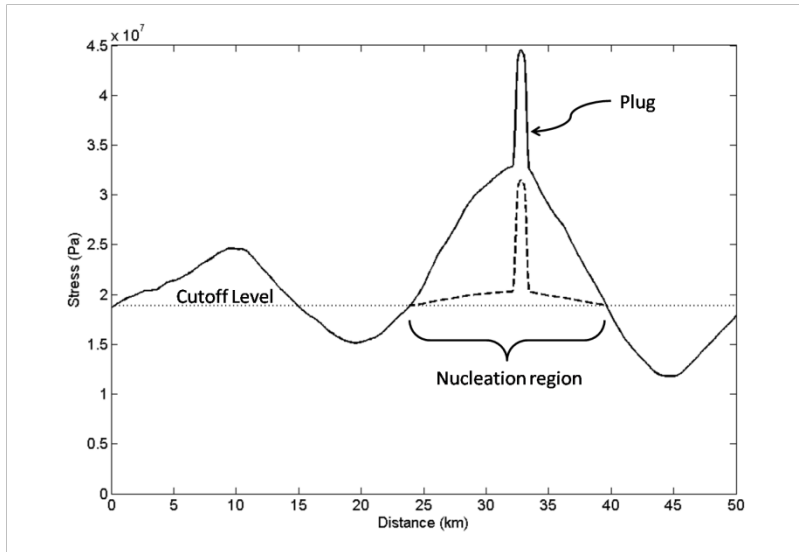


Figure 4.4: Illustration of scaling the nucleation stress. The point of highest shear stress in the model is identified, and a cosine-shaped plug of radius R_{plug} and amplitude τ_{plug} is superimposed. A contour is taken around the plug at a set cutoff stress (in the 2D modeling this is halfway between the maximum and minimum stress, in the 3D modeling it is 65% of the way to the maximum stress). The stress in the "nucleation region" within this contour is then scaled in 10 equal increments of 10% down to the cutoff stress level. The dashed line shows the lowest stress step and the solid line shows the highest stress step, equal to the original stress level.

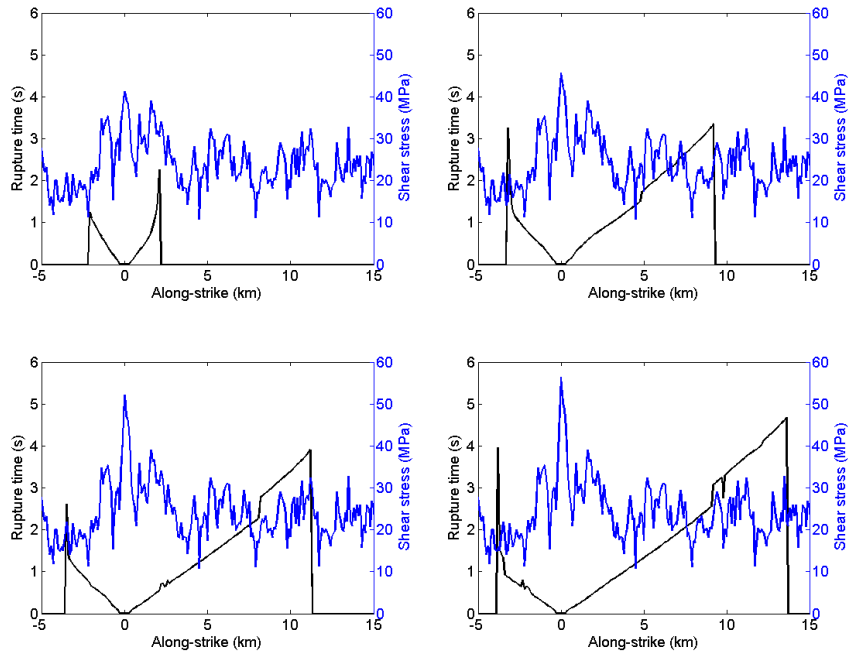


Figure 4.5: Example of model exhibiting the hypothesized behavior. Black trace shows the time of rupture with units on the left axis, and blue trace shows the initial shear stress with units on the right axis. At low nucleation stress an area of only ~ 4 km has ruptured (upper left), and with progressively higher nucleation stress more of the fault ruptures until, at the full value of nucleation stress, ~ 16 km of the fault has ruptured.

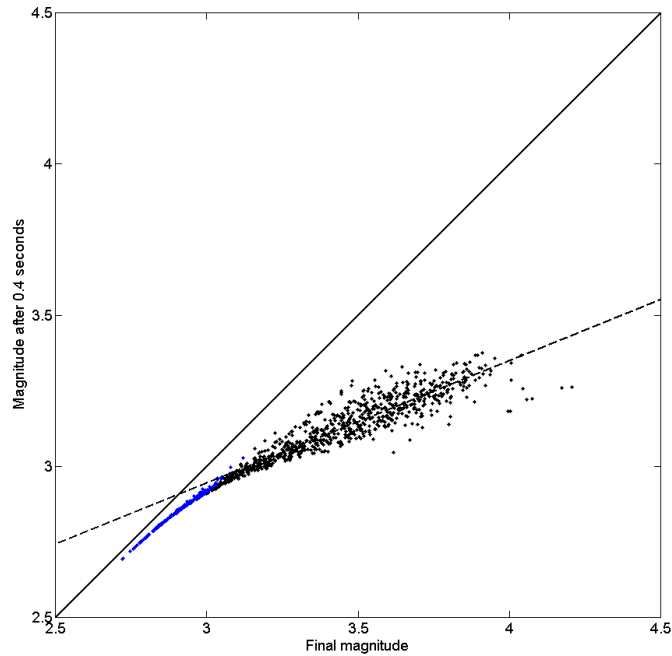


Figure 4.6: Final magnitude of 100 realizations of the best individual in the 2D search, plotted against initial magnitude (defined as magnitude after 0.4 seconds of rupture). Blue points have an initial magnitude greater than 97% of the final magnitude, and are excluded from the regression. The best-fit slope using only the black points is greater than zero at 95% confidence, indicating widespread dependence of final magnitude on the early rupture in these models.

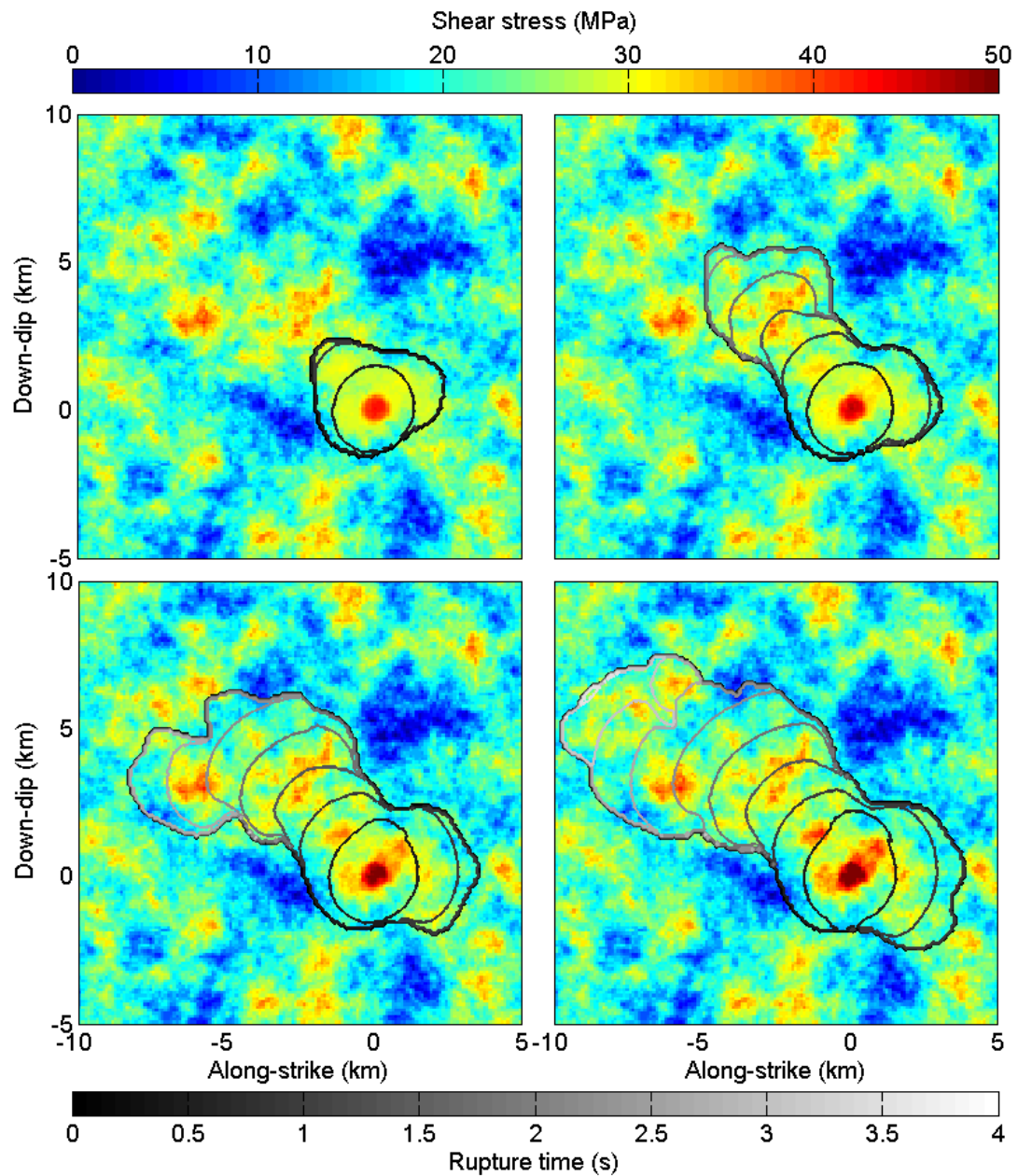


Figure 4.7: Example of a 3D model exhibiting the hypothesized behavior. The color field represents the initial shear stress, and the contours represent rupture time in 0.2 second increments. When nucleation stress is scaled down, only the nucleation region ruptures (upper left). As the nucleation stress is scaled up progressively, a second, third and finally fourth asperity ruptures when the stress is at its full value.

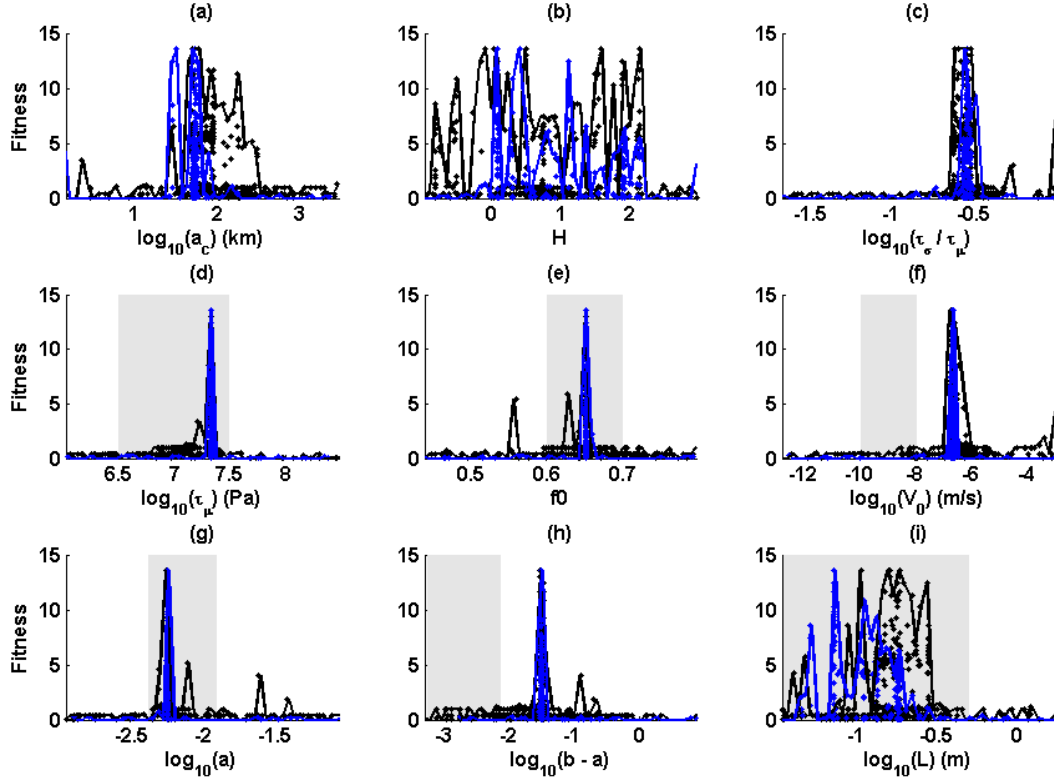


Figure 4.8: The fitness of all the individuals in the 2D search (black points, 630 individuals) and the 3D search (blue points, 180 individuals) as a function of each of 4 stress-related parameters (a-d) and each of 5 friction-related parameters (e-i). Maximum fitness envelopes are plotted as black lines (2D search) and blue lines (3D search). Where previous work constrains the possible values of the parameters, the range of possible values is shown as a light shading (d-i).

5. Conclusion

We now have two lines of evidence suggesting that earthquakes are not purely cascading phenomena, and that the early history of rupture exerts a degree of influence on the final magnitude of an event. The first, the observation that the amplitude and spectral content of P-waves scale with event magnitude, is in agreement with prior studies [*Allen and Kanamori, 2003; Olson and Allen, 2005; Lockman and Allen, 2007; Lewis and Ben-Zion, 2008*] that show this behavior at a variety of scales and in different tectonic contexts. The second line of evidence, a statistical observation that the early slip history of kinematic inversions scales with event magnitude, is an independent study that corroborates the seismic observations and supports deterministic behavior, at least in aggregate statistics, with a high degree of confidence.

With these two lines in hand, we test the hypothesis that this aggregate determinism is driven by variations in the intensity of the early rupture phase, which impart more or less energy to the rupture front and thus affect its ability to overcome barriers further along the rupture. We simulate ruptures with dynamic models of 2D and 3D spaces, and search for the hypothesized behavior using a Genetic Algorithm. We find that the hypothesized behavior occurs under realistic conditions for the most part, though it requires a greater degree of velocity-weakening behavior than is generally observed in lab friction experiments. This velocity-weakening behavior may be generated in the field by flash heating of fault materials and thermal pressurization of pore fluids, but these effects need to be accounted for in the modeling before definitively supporting the hypothesis.

We are now a few steps closer to answering the question of whether earthquakes are deterministic or cascading phenomena. However, much work remains to be done in this field. In terms of the observational data, a great deal must be done to improve the statistical significance of the observed scaling, especially at large magnitudes. This is difficult because of the small number of large earthquakes that occurs each year, few of which occur in the vicinity of high-quality regional seismic networks. As far as modeling the physics of the observed scaling, many questions remain unanswered, such as the specific nature of the “intensity” of the early rupture phase and what determines this intensity. Finally, whatever specific physical mechanism is hypothesized to drive the determinism of earthquake ruptures, this mechanism must be able to generate the observed effects on the properties of P-waves.

The extant questions about the physics of earthquake rupture are no longer merely academic exercises. The advent of P-wave-based Earthquake Early Warning systems makes earthquake source studies into an eminently practical field, with significant implications regarding the applicability of these warnings to large events with long, complex rupture histories. We are rapidly approaching the point where a thorough understanding of the physics of rupture will make it possible to eke out a few precious extra seconds of

warning, or to converge on ground motion estimates with greater confidence and smaller errors. The answers to questions of source physics may effect greater savings of life and property in the near future. There is no better time to study these processes, because today these questions are perhaps more important than any others in modern seismology.

6. References

- Abercrombie, R. (2005). Seismology - The start of something big?, *Nature* **439** 171-173, doi:10.1038/438171a.
- Abercrombie, R., and P. Leary (1993). Source parameters of small earthquakes recorded at 2.5 km depth, Cajon Pass, Southern California, *Geophys. Res. Lett.* **20** 1511-1514.
- Abercrombie, R., and J. Mori (1994). Local observations of the onset of a large earthquake - 25 June 1992 Landers, California, *Bull. Seism. Soc. Am.* **94** 725-734.
- Aki, K. (2000). Scale-dependence in earthquake processes and seismogenic structures, *Pure Appl. Geophys.* **157** 2249-2258.
- Allen, R. V. (1978). Automatic earthquake recognition and timing from single traces, *Bull. Seism. Soc. Am.* **68** 1520-1532.
- Allen, R. M., (2004). Rapid magnitude determination for earthquake early warning, in M. Pecce, G. Manfredi and A. Zollo (Eds.) *The many facets of seismic risk*. Universita degli Studi di Napoli "Federico II", Napoli, Italy.
- Allen, R. M. (2006). Probabilistic warning times for earthquake ground shaking in the San Francisco Bay Area, *Seism. Res. Lett.* **77** 371-376.
- Allen, R. M., (2007). The ElarmS earthquake early warning methodology and its application across California, in P. Gasparini, G. Manfredi and J. Zschau (Eds.) *Earthquake Early Warning Systems*. Springer, pp.21-44.
- Allen, R. M., and H. Kanamori (2003). The potential for earthquake early warning in southern California, *Science* **300** 786-789.
- Aochi, H., and S. Ide (2004). Numerical study on multi-scaling earthquake rupture, *Geophys. Res. Lett.* **31** L02606, doi:10.1029/2003GL018708.
- Bak, P., and C. Tang (1989). Earthquakes as a self-organized critical phenomenon, *J. Geophys. Res.* **94** 15635-15637.
- Beaty, K. S., D. R. Schmitt, and M. Sacchi (2002). Simulated annealing inversion of multimode Rayleigh wave dispersion curves for geological structure, *Geophys. J. Int.* **151** 622-631.
- Beresnev, I. A. (2003). Uncertainties in finite-fault slip inversions: To what extent to believe? (A critical review), *Bull. Seism. Soc. Am.* **93** 2445-2458.
- Beroza, G. C., and W. L. Ellsworth (1996). Properties of the seismic nucleation phase, *Tectonophysics* **261** 209-227.
- Blanpied, M. L., D. A. Lockner, and J. D. Byerlee (1991). Fault stability inferred from granite sliding experiments at hydrothermal conditions, *Geophys. Res. Lett.* **18** 609-612.
- Boatwright, J., H. Bundock, J. Luetgert, L. Seekins, L. Gee, and P. Lombard (2003). The dependence of PGA and PGV on distance and

magnitude inferred from northern California ShakeMap data, *Bull. Seism. Soc. Am.* **68** 2043-2055.

Boore, D. M., W. B. Joyner, and T. E. Fumal (1997). Equations for estimating horizontal response spectra and peak accelerations from western North American earthquakes: A summary of recent work, *Seism. Res. Lett.* **68** 128-153.

Borcherdt, R. D. (1994). Estimates of site-dependent response spectra for design (methodology and justification), *Earthquake Spectra* **10** 617-654.

Brown, H., R. M. Allen, and V. F. Grasso (2009). Testing ElarmS in Japan, *Seism. Res. Lett.* **80** 727-739, doi:10.1785/gssrl.80.5.727.

Choy, G. L. (1995). Global patterns of radiated seismic energy and apparent stress, *J. Geophys. Res.* **100** 18205-18228.

Cotton, F., and M. Campillo (1995). Frequency-domain inversion of strong motions - Application to the 1992 Landers earthquake, *J. Geophys. Res.* **100** 3961-3975.

Dieterich, J. H. (1978). Time-dependent friction and the mechanics of stick-slip, *Pure Appl. Geophys.* **116** 790-806.

Dieterich, J. H., and B. D. Kilgore (1994). Direct observation of frictional contacts: new insights for state-dependent properties, *Pure Appl. Geophys.* **143** 283-302.

Dieterich, J. H., and B. D. Kilgore (1996). Imaging surface contacts: power law contact distributions and contact stresses in quartz, calcite, glass and acrylic plastic, *Tectonophysics* **2556** 219-239.

Dunham, E. M. (2005). Dissipative interface waves and the transient response of a three dimensional sliding interface with Coulomb friction, *J. Mech. Phys. Solids* **53** 327-357, doi:10.1016/j.jmps.2004.07.003.

Ellsworth, W. L., and G. C. Beroza (1995). Seismic evidence for an earthquake nucleation phase, *Science* **268** 851-855.

Ellsworth, W. L., and G. C. Beroza (1998). Observation of the seismic nucleation phase in the Ridgecrest, California earthquake sequence, *Geophys. Res. Lett.* **25** 401-404.

Espinosa Aranda, J. M., A. Jimenez, G. Ibarrola, F. Alcantar, A. Aguilar, M. Inostroza, and S. Maldonado (1995). Mexico City seismic alert system, *Seism. Res. Lett.* **66** 42-53.

Federal Emergency Management Agency, (2008). *HAZUS® MH Estimated Annualized Earthquake Losses for the United States*, FEMA 336

Fukuyama, E., C. Hashimoto, and M. Matsu'ura (2002). Simulation of the transition of earthquake rupture from quasi-static growth to dynamic propagation, *Pure Appl. Geoph.* **179** 2057-2066.

Fukuyama, E., and R. Madariaga (2000). Dynamic propagation and interaction of a rupture front on a planar fault, *Pure Appl. Geoph.* **157** 1959-1979.

Hartzell, S. H., and T. H. Heaton (1983). Inversion of strong ground motion and teleseismic waveform data for the fault rupture history of the

1979 Imperial Valley, California, earthquake, *Bull. Seism. Soc. Am.* **73** 1553-1583.

Haskell, N. A. (1964). Total energy and energy spectral density of elastic wave radiation from propagating faults, *Bull. Seism. Soc. Am.* **54** 1811-1841.

Heaton, T. H. (1990). Evidence for and implications of self-healing pulses of slip in earthquake rupture, *Phys. Earth Planet. Inter.* **64** 1-20.

Hickman, S. M., and M. D. Zoback (2004). Stress orientations and magnitudes in the SAFOD pilot hole, *Geophys. Res. Lett.* **31** L15S12, doi:10.1029/2004GL020043.

Hill, D. P., J. P. Eaton, and L. M. Jones, (1990). Seismicity, 1980-86, in R.E. Wallace (Ed.) *The San Andreas Fault System*. U.S. Geological Survey Professional Paper 1515.

Ide, S., and H. Aochi (2005). Earthquakes as multiscale dynamic ruptures with heterogeneous fracture surface energy, *J. Geophys. Res.* **110** B11303, doi:10.1029/2004JB003591.

Ide, S., and G. C. Beroza (2001). Does apparent stress vary with earthquake size?, *Geophys. Res. Lett.* **28** 3349-3352.

Ide, S., and M. Takeo (1997). Determination of constitutive relations of fault slip based on seismic wave analysis, *J. Geophys. Res.* **102** 27379-27391.

Iio, Y. (1992). Slow initial phase of the P-wave velocity pulse generated by microearthquakes, *Geophys. Res. Lett.* **19** 477-480.

Iio, Y. (1995). Observations of the slow initial phase generated by microearthquakes - implications for earthquake nucleation and propagation, *J. Geophys. Res.* **100** 15333-15349.

Ji, C., D. J. Wald, and D. V. Helmberger (2002). Source description of the 1999 Hector Mine, California, earthquake, Part I: Wavelet domain inversion theory and resolution analysis, *Bull. Seism. Soc. Am.* **92** 1192-1207.

Kaverina, A., D. S. Dreger, and E. Price (2002). The combined inversion of seismic and geodetic data for the source process of the 16 October 1999 M-w 7.1 Hector Mine, California earthquake, *Bull. Seism. Soc. Am.* **92** 1266-1280.

Kilb, D., and J. Gomberg (1999). The initial subevent of the 1994 Northridge, California, earthquake: Is earthquake size predictable?, *J. Seismol.* **3** 409-420.

Kircher, C. A., H. A. Seligson, J. Bouabid, and G. C. Morrow (2006). When the Big One strikes again - Estimated losses due to a repeat of the 1906 San Francisco earthquake, *Earthquake Spectra* **22** S297-S339, doi:10.1193/1.2187067.

Lapusta, N., and J. R. Rice (2003). Nucleation and early seismic propagation of small and large events in a crustal earthquake model, *J. Geophys. Res.* **108** 2205, doi:10.1029/2001JB000793.

Lewis, M. A., and Y. Ben-Zion (2008). Examination of scaling between earthquake magnitude and proposed early signals in P waveforms from very near source stations in a South African gold mine, *J. Geophys. Res.* **113** B09305, doi:10.1029/2007JB005506.

Liu, P., S. Hartzell, and W. Stephenson (1995). Non-linear multiparameter inversion using a hybrid global search algorithm: applications in reflection seismology, *Geophys. J. Int.* **122** 991-1000.

Lockman, A. B., and R. M. Allen (2005). Single-station earthquake characterization for early warning, *Bull. Seism. Soc. Am.* **95** 2029-2039.

Lockman, A. B., and R. M. Allen (2007). Magnitude-period scaling relations for Japan and the Pacific Northwest: Implications for earthquake early warning, *Bull. Seism. Soc. Am.* **97** 140-150.

Loveless, J. P., M. E. Pritchard, and N. Kukowski (in review). Testing mechanisms of seismic segmentation with slip distributions from recent earthquakes along the Andean margin, *Tectonophysics*.

Mai, P. M., P. Spudich, and J. Boatwright (2005). Hypocenter locations in finite-source rupture models, *Bull. Seism. Soc. Am.* **95** 965-980.

Marone, C. (1998). Laboratory-derived friction laws and their application to seismic faulting, *Ann. Rev. Earth Planet. Sci.* **26** 643-696.

Mayeda, K., and W. R. Walter (1996). Moment, energy, stress drop, and source spectra of western United States earthquakes from regional coda envelopes, *J. Geophys. Res.* **101** 11195-11208.

Mikumo, T., K. B. Olsen, E. Fukuyama, and Y. Yagi (2003). Stress-breakdown time and slip-weakening distance inferred from slip-velocity functions on earthquake faults, *Bull. Seism. Soc. Am.* **93** 264-282.

Mori, J., and H. Kanamori (1996). Initial rupture of earthquakes in the 1995 Ridgecrest, California sequence, *Geophys. Res. Lett.* **23** 2437-2440.

Murphy, S., and S. Nielsen (2009). Estimating earthquake magnitude with early arrivals: a test using dynamic and kinematic models, *Bull. Seism. Soc. Am.* **99** 1-23.

Nakatani, M., S. Kaneshima, and Y. Fukao (2000). Size-dependent microearthquake initiation inferred from high-gain and low-noise observations at Nikko district, Japan, *J. Geophys. Res.* **105** 28095-28109.

Newmark, N. M., and W. J. Hall (1982). Earthquake spectra and design, *Geotechnique* **25** 139-160.

Noda, H., E. M. Dunham, and J. R. Rice (2009). Earthquake ruptures with thermal weakening and the operation of major faults at low overall stress levels, *J. Geophys. Res.* **114** B07302, doi:10.1029/2008JB006143.

Odaka, T., S. Ashiya, S. Tsukada, S. Sato, K. Ohtake, and D. Nozaka (2003). A new method of quickly estimating epicentral distance and magnitude from a single seismic record, *Bull. Seism. Soc. Am.* **93** 526-532.

Oglesby, D. D., and S. M. Day (2002). Stochastic fault stress: Implications for fault dynamics and ground motion, *Bull. Seism. Soc. Am.* **92** 3006-3021.

Oglesby, D. D., D. S. Dreger, R. A. Harris, N. Ratchkovski, and R. Hansen (2004). Inverse kinematic and forward dynamic models of the 2002 Denali fault earthquake, Alaska, *Bull. Seism. Soc. Am.* **94** S214-S233.

Ohnaka, M. (2000). A physical scaling relation between the size of an earthquake and its nucleation zone size, *Pure Appl. Geoph.* **157** 2259-2282.

Ohnaka, M. (2004). A constitutive scaling law for shear rupture that is inherently scale-dependent, and physical scaling of nucleation time to critical point, *Pure Appl. Geoph.* **161** 1915-1929.

Ohnaka, M., and L.-f. Shen (1999). Scaling of the shear rupture process from nucleation to dynamic propagation: Implications of geometric irregularity of the rupturing surfaces, *J. Geophys. Res.* **104** 817-844.

Olson, E. L., and R. M. Allen (2005). The deterministic nature of earthquake rupture, *Nature* **438** 212-215, doi:10.1038/nature04214.

Olson, A. H., and R. J. Apsel (1982). Finite faults and inverse-theory with applications to the 1979 Imperial-Valley earthquake, *Bull. Seism. Soc. Am.* **72** 1969-2001.

Otsuki, K., and T. Dilov (2005). Evolution of hierarchical self-similar geometry of experimental fault zones: Implications for seismic nucleation and earthquake size, *J. Geophys. Res.* **110** B03303, doi:10.1029/2004JB003359.

Pasyanos, M. E., D. S. Dreger, and B. Romanowicz (1996). Toward real-time estimation of regional moment tensors, *Bull. Seism. Soc. Am.* **86** 1255-1269.

Petersen, M. D., A. D. Frankel, S. C. Harmsen, C. S. Mueller, K. M. Haller, R. L. Wheeler, R. L. Wesson, Y. Zeng, O. S. Boyde, D. M. Perkins, N. Luco, F. E. H, C. J. Wills, and K. S. Rukstales, (2008). *Documentation for the 2008 Update of the United States National Seismic Hazard Maps*, U.S. Geological Survey Open-File Report 2008-1128

Pritchard, M. E., and E. J. Fielding (2008). A study of the 2006 and 2007 earthquake sequence of Pisco, Peru, with InSAR and teleseismic data, *Geophys. Res. Lett.* **35** L09308, doi:10.1029/2008GL033374.

Pritchard, M. E., C. Ji, and M. Simons (2006). Distribution of slip from 11 Mw > 6 earthquakes in the northern Chile subduction zone, *J. Geophys. Res.* **111** B10302, doi:10.1029/2005JB004013.

Pritchard, M. E., E. O. Norabuena, C. Ji, R. Boroschek, D. Comte, M. Simons, T. Dixon, and P. A. Rosen (2007). Geodetic, teleseismic, and strong motion constraints on slip from recent southern Peru subduction zone earthquakes, *J. Geophys. Res.* **112** B03307, doi:10.1029/2006JB004294.

Rice, J. R. (1983). Constitutive relations for fault slip and earthquake instabilities, *Pure Appl. Geophys.* **121** 443-475.

Rice, J. R., N. Lapusta, and K. Ranjith (2001). Rate and state dependent friction and the stability of sliding between elastically deformable solids, *J. Mech. Phys. Solids* **49** 1865-1898.

- Ripperger, J., J.-P. Ampuero, P. M. Mai, and D. Giardini (2007). Earthquake source characteristics from dynamic rupture with constrained stochastic fault stress, *J. Geophys. Res.* **112** B04311, doi:10.1029/2006JB004515.
- Ruina, A. (1983). Slip instability and state variable friction laws, *J. Geophys. Res.* **88** 10359-10370.
- Rydelek, P., and S. Horiuchi (2006). Earth science: Is earthquake rupture deterministic?, *Nature* **442** E5-E6.
- Sambridge, M., and G. Drijkoningen (1992). Genetic algorithms in seismic waveform inversion, *Geophys. J. Int.* **109** 323-342.
- Satake, K. (1987). Inversion of tsunami waveforms for the estimation of a fault heterogeneity: Method and numerical experiments, *J. Phys. Earth* **35** 241-254.
- Sato, T., and H. Kanamori (1999). Beginning of earthquakes modeled with the Griffith's fracture criterion, *Bull. Seism. Soc. Am.* **89** 80-93.
- Sato, K., and J. Mori (2006). Scaling relationship of initiations for moderate to large earthquakes, *J. Geophys. Res.* **111** B05306, doi:10.1029/2005JB003613.
- Scholz, C. H. (1988). The critical slip distance for seismic faulting, *Nature* **336** 761-763.
- Scholz, C. H. (1998). Earthquakes and friction laws, *Nature* **391** 37-42.
- Sekiguchi, H., K. Irikura, and T. Iwata (2000). Fault geometry at the rupture termination of the 1995 Hyogo-ken Nanbu earthquake, *Bull. Seism. Soc. Am.* **90** 117-133.
- Sleeman, R., and T. van Eck (1999). Robust automatic P-phase picking: an on-line implementation in the analysis of broadband seismogram recordings, *Phys. Earth Planet. Inter.* **113** 265-275.
- Somerville, P. G., K. Irikura, R. Graves, S. Sawada, D. Wald, N. Abrahamson, Y. Iwasaki, T. Kagawa, N. Smith, and A. Kowada (1999). Characterizing crustal earthquake slip models for the prediction of strong ground motion, *Seism. Res. Lett.* **70** 59-80.
- Steacy, S. J., and J. McCloskey (1998). What controls an earthquake's size? Results from a heterogeneous cellular automaton, *Geophys. J. Int.* **133** F11-F14.
- Stoffa, P. L., and M. K. Sen (1991). Nonlinear multiparameter optimization using genetic algorithms: Inversion of plane-wave seismograms, *Geophysics* **56** 1794-1810.
- Tanioka, Y., and K. Satake (2001). Coseismic slip distribution of the 1946 Nankai earthquake and aseismic slips caused by the earthquake, *Earth Planets and Space* **53** 235-241.
- Tembe, S., D. Lockner, and T.-f. Wong (2010). Constraints on the stress state of the San Andreas Fault with analysis based on core and cuttings from San Andreas Fault Observatory at Depth (SAFOD) drilling phases 1 and 2, *J. Geophys. Res.* **115** B03418, doi:10.1029/2009JB000818.

Tullis, T. E., and J. D. Weeks (1986). Constitutive behavior and stability of frictional sliding of granite, *Pure Appl. Geoph.* **124** 363-414.

Umeda, Y. (1990). High-amplitude seismic waves radiated from the bright spot of an earthquake, *Tectonophysics* **175** 81-92.

Umeda, Y. (1992). The bright spot of an earthquake, *Tectonophysics* **211** 13-22.

Vera, E. E., J. C. Mutter, P. Buhl, J. A. Orcutt, A. J. Harding, M. E. Kappus, R. S. Detrick, and T. M. Brocher (1990). The structure of 0- to 0.2-m.y.-old oceanic crust at 9 degrees N on the East Pacific Rise from expanded spread profiles, *J. Geophys. res.* **95** 15529-15556.

Wald, D. J., T. H. Heaton, and K. W. Hudnut (1996). The slip history of the 1994 Northridge, California, earthquake determined from strong-motion, teleseismic, GPS, and leveling data, *Bull. Seism. Soc. Am.* **86** S49-S70.

Wald, D. J., V. Quitoriano, T. H. Heaton, and H. Kanamori (1999b). Relationship between peak ground acceleration, peak ground velocity, and Modified Mercalli Intensity for earthquakes in California, *Earthquake spectra* **15** 557-564.

Wald, D. J., V. Quitoriano, T. H. Heaton, H. Kanamori, C. W. Scrivner, and B. C. Worden (1999a). TriNet ShakeMaps: Rapid generation of instrumental ground motion and intensity maps for earthquakes in southern California, *Earthquake Spectra* **15** 537-556.

Wald, D. J., B. C. Worden, V. Quitoriano, and K. L. Pankow, (2005). *ShakeMap® manual: Technical manual, users guide, and software guide*, U.S. Geological Survey Techniques and Methods 12-A1

Wells, D. L., and K. J. Coppersmith (1994). New empirical relationships among magnitude, rupture length, rupture width, rupture area, and surface displacement, *Bull. Seism. Soc. Am.* **84** 974-1002.

Wibberley, C. A.J., and T. Shimamoto (2005). Earthquake slip weakening and asperities explained by thermal pressurization, *Nature* **436** 689-692, doi:10.1038/nature03901.

Wills, C. J., M. D. Petersen, W. A. Bryant, M. S. Reichle, G. J. Saucedo, S. S. Tan, G. C. Taylor, and J. A. Treiman (2000). A site-conditions map for California based on geology and shear wave velocity, *Bull. Seism. Soc. Am.* **90** S187-S208.

Wolfe, C. J. (2006). On the properties of predominant-period estimators for earthquake early warning, *Bull. Seism. Soc. Am.* **96** 1961-1965.

Working Group on California Earthquake Probabilities, (2008). *The Uniform California Earthquake Rupture Forecast, Version 2 (UCERF2)*, U.S. Geological Survey Open File Report 2007-1437

Wu, Y.-M., and H. Kanamori (2005a). Experiment on an onsite early warning method for the Taiwan early warning system, *Bull. Seism. Soc. Am.* **95** 347-353.

Wu, Y.-M., and H. Kanamori (2005b). Rapid assessment of damage potential of earthquakes in Taiwan from the beginning of P waves, *Bull. Seism. Soc. Am.* **95** 1181-1185.

Wurman, G., R. M. Allen, D. S. Dreger, and D. D. Oglesby (in review). Statistical Testing of Theoretical Rupture Models Against Kinematic Inversions, *Bull. Seism. Soc. Am.*

Wurman, G., R. M. Allen, and P. Lombard (2007). Toward earthquake early warning in northern California, *J. Geophys. Res.* **112** B08311, doi:10.1029/2006JB004830.

Wu, Y.-M., H.-Y. Yen, L. Zhao, B.-S. Huang, and W.-T. Liang (2006). Magnitude determination using initial P waves: A single-station approach, *Geophys. Res. Lett.* **33** L05306, doi:10.1029/2005GL025395.

Wyss, M., and J. N. Brune (1967). Alaska earthquake of 28 March 1964 - A complex multiple rupture, *Bull. Seism. Soc. Am.* **57** 1017-1023.

Yoshida, S., and K. Koketsu (1990). Simultaneous inversion of wave-form and geodetic data for the rupture process of the 1984 Naganoken Seibu, Japan, earthquake, *Geophys. J. Int.* **103** 355-362.

Zhang, H. J., C. Thurber, and C. Rowe (2003). Automatic P-wave arrival detection and picking with multiscale wavelet analysis for single-component recordings, *Bull. Seism. Soc. Am.* **93** 1904-1912.