

CSDL-T-1214

**LIMITING VIBRATION IN SYSTEMS
WITH CONSTANT AMPLITUDE ACTUATORS
THROUGH COMMAND PRESHAPING**

by
Keith Eric Rogers

May 1994

**Master of Science Thesis
Massachusetts Institute of Technology**

(NASA-CR-188284) LIMITING
VIBRATION IN SYSTEMS WITH CONSTANT
AMPLITUDE ACTUATORS THROUGH COMMAND
PRESHAPING M.S Thesis - MIT
(Draper (Charles Stark) Lab.)
117 p

N94-32907

Unclass

G3/39 0010459



The Charles Stark Draper Laboratory, Inc.
555 Technology Square, Cambridge, Massachusetts 02139-3563

Limiting Vibration in Systems with Constant Amplitude Actuators through Command Preshaping

by

Keith Eric Rogers

B.S.E., Princeton University
(June, 1992)

SUBMITTED TO THE DEPARTMENT OF AERONAUTICS AND
ASTRONAUTICS IN PARTIAL FULFILLMENT OF THE
REQUIREMENTS FOR THE DEGREE OF

MASTER OF SCIENCE

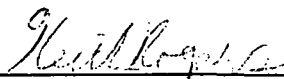
at the

MASSACHUSETTS INSTITUTE OF TECHNOLOGY

May 1994

© 1994 Keith Eric Rogers
All Rights Reserved

Signature of Author



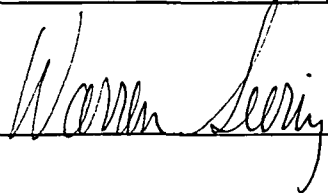
Department of Aeronautics and Astronautics
May, 1994

Approved by



Technical Supervisor, Draper Laboratory
Kevin Gift

Certified by



Professor Warren P. Seering
Thesis Supervisor

Accepted by

Professor Harold Y. Wachman
Chairman, Departmental Graduate Committee

Limiting Vibration in Systems with Constant Amplitude Actuators through Command Preshaping

by

Keith Eric Rogers

B.S.E., Princeton University
(June, 1992)

Submitted to the Department of Aeronautics and
Astronautics on May 6, 1994 in partial fulfillment of the
requirements for the degree of Master of Science

Abstract

The basic concepts of command preshaping were taken and adapted to the framework of systems with constant amplitude (on-off) actuators. In this context, pulse sequences were developed which help to attenuate vibration in flexible systems with high robustness to errors in frequency identification. Sequences containing impulses of different magnitudes were approximated by sequences containing pulses of different durations. The effects of variation in pulse width on this approximation were examined. Sequences capable of minimizing loads induced in flexible systems during execution of commands were also investigated. The usefulness of these techniques in real-world situations was verified by application to a high fidelity simulation of the space shuttle. Results showed that constant amplitude preshaping techniques offer a substantial improvement in vibration reduction over both the standard and upgraded shuttle control methods and may be mission enabling for use of the shuttle with extremely massive payloads.

Thesis Supervisor:
Warren P. Seering

PRECEDING PAGE BLANK NOT FILMED

Acknowledgment

May 3, 1994

This thesis was prepared at the Charles Stark Draper Laboratory, Inc., under contract NAS9-18426.

Publication of this thesis does not constitute approval by Draper or the sponsoring agency of the findings or conclusions contained herein. It is published for the exchange and stimulation of ideas.

I hereby assign my copyright of this thesis to the Charles Stark Draper Laboratory, Inc., Cambridge, Massachusetts.

Neil Rogers

Permission is hereby granted by the Charles Stark Draper Laboratory, Inc., to the Massachusetts Institute of Technology to reproduce any or all of this thesis.

PRECEDING PAGE BLANK NOT FILMED

RECEIVED 1964 APR 10 10 10 AM

Acknowledgments

There are quite a few people who helped me get to this point, and I would like to take this space to thank them.

Thanks first to the Charles Stark Draper Laboratory for providing me with the financial and academic support which have allowed me to complete my Masters Degree at MIT.

Next, I'd like to thank my supervisor at Draper, Kevin Gift, and my advisor at MIT, Professor Warren Seering. They have both done a great job of keeping me on track and providing with feedback on my work.

Thanks also to my officemates at Draper and the AI lab, Mark Jackson and Andy Christian, who were both excellent sources of ideas and humored me when I needed to bounce my own off of someone else. Andy, your experience with preshaping and insight into the issues involved were a great, great help.

I'd also like to express my appreciation of all the folks in the (former) Manned Space Flight Group at Draper who were always willing to lend a hand and answer questions about the shuttle.

Finally, I'd like to thank my mother and brother for all the love and support that they have given me over the years. Without them, I could not have gotten here.

~~PRECEDING~~ PAGE BLANK NOT FILMED

PAGE 6 INTENTIONALLY BLANK

Table of Contents

Acknowledgments.....	7
Table of Contents.....	9
List of Figures.....	11
1. Introduction.....	13
2. Shuttle Operations.....	19
2.1 The Digital AutoPilot.....	20
2.2 The Remote Manipulator System.....	23
2.3 Maneuvers.....	24
2.4 Concerns with Dynamic Interactions.....	26
2.5 Current Solutions.....	29
2.6 The Role of Command Preshaping.....	31
3. Impulse Preshaping in Constant Amplitude Systems.....	33
3.1 Theoretical Model.....	35
3.2 The Single Pulse.....	37
3.3 The Line Sequence.....	40
3.3.1 Time Cost.....	42
3.4 The L^2 sequence.....	44
3.4.1 Concatenation.....	47
3.4.2 Pulse Width.....	48
3.4.3 Time Cost.....	54
3.4.4 Resolution.....	55
3.5 Y and S Sequences.....	57
3.6 Load Sequences.....	63
3.7 More Complex Systems.....	69
3.8 Designing a Constant Amplitude Preshaper.....	75
4. Application to the Shuttle.....	79
4.1 The DRS.....	80
4.2 DRS Configuration.....	82
4.3 Frequency Identification.....	84
4.4 Simulation Results.....	89
3.4.1 The Upper Atmospheric Research Satellite.....	90
3.4.2 The Hubble Space Telescope.....	99
3.4.3 Space Station Model MB6.....	107
3.4.4 Unloaded Arm.....	112
5. Conclusion.....	115
5.1 Future Work.....	118
Bibliography.....	120

List of Figures

2.1	Locations and directions of the shuttle's primary and vernier jets.....	22
2.2	Phase plane deadbands for a single axis.....	23
2.3	Orbiter attitude in two phases of deployment operations for the Gamma Ray Observatory.....	25
3.1	Two-mass, one spring model used for analysis in this chapter.....	35
3.2	Typical thrust profile for a space shuttle Primary Rotation Control System (PRCS) jet firing a minimum duration (80 millisecond) pulse.....	36
3.3	The single pulse.....	38
3.4	Graph of the residual amplitude from a single pulse normalized by the command size vs. the command size.....	39
3.5	The line sequence.....	41
3.6	Force Profile for a line sequence requiring a burn time longer than one full period.....	43
3.7	Time cost as a function of command size.....	44
3.8	The L^2 sequence.....	45
3.8	Construction of large sequences through concatenation.....	48
3.10	Residual magnitude at the design frequency for an L^2 sequence with varying pulse width.....	50
3.11	Residual magnitude for an L^2 sequence as a function of frequency and pulse width.....	51
3.12	Concatenation of an L^2 sequence for large command sizes.....	55
3.13	Force profiles for different methods of splitting up a total firing time of 19 times the minimum pulse width.....	57
3.14	The Y and S sequences.....	58
3.15	Insensitivity curves for L, Y, and S sequences.....	60
3.16	Insensitivity curves for the L^3 sequence and the LY@2/3 sequence.....	61
3.17	The insensitivity curve of an LY sequence.....	62
3.18	Load sequence for $n = 1$	66
3.19	Map of the effect of sequence length on insensitivity for load sequences.....	68
3.20	Effect of damping ratio on impulse magnitude for an L^2 sequence.....	72
3.21	Error induced by the assumption of zero damping as a function of actual damping ratio for three different sequences.....	73
3.22	Illustration of the effects of convolution on sequence insensitivity.....	74
4.1	Comparison of results from a DRS simulation with actual flight data from STS-61, the Hubble Space Telescope servicing mission.....	83
4.2	Power spectral densities of RSS POR position from simulation of PRCS jet firings in two different directions.....	87
4.3	POR RSS time history for simulation output from 0.48 second -Roll, -Pitch PRCS jet firing.....	88
4.4	The Upper Atmospheric Research Satellite deployed on the RMS in capture/release configuration.....	91

4.5	Simulation results from primary jet firings of total duration equal to 0.48 seconds.....	93
4.6	A test of LY sequence robustness.....	94
4.7	Close-up of the time history shown in Figure 4.5.....	96
4.8	Time histories for VRCS jet firings totaling 9.6 seconds.....	98
4.9	The Hubble Space Telescope deployed on the RMS in extended park position.....	99
4.10	Power Spectral Density of POR response to a 0.08 second -Pitch, -Roll PRCS jet firing.....	102
4.11	Time histories for PRCS jet firings totaling 0.54 seconds.....	103
4.12	Time histories for VRCS jet firings totaling 9.6 seconds.....	104
4.13	The insensitivity curve of an $L^2@3/2Y$ sequence.....	105
4.14	A look at the effect of pulse size on preshaping sequence performance.....	106
4.15	The MB6 space station configuration with shuttle attached via the RMS as planned in 1992.....	108
4.16	PSD for MB6 and matching sequence insensitivity curve.....	110
4.17	Time histories for VRCS jet firings totaling 38.4 seconds.....	111
4.18	Time histories for PRCS jet firings totaling 0.64 seconds.....	112
4.19	Time histories for PRCS jet firings totaling 0.48 seconds for an unloaded arm.....	113

Introduction

Chapter 1

Over the past seven or eight years a cooperative effort between Draper Lab and MIT has produced quite a bit of work in the field of command preshaping. Also known as impulse shaping or input shaping, this new area of research offers substantial promise in the search for control methods which can limit vibration in flexible systems. By convolving sequences of impulses tuned to the system's modes of vibration with commands, preshapers are able to reduce residual vibration in complex flexible systems. In addition, they offer robustness to uncertainty in system frequencies, a minimal time delay formulation, and the ability to handle multiple frequencies of vibration.

So far, however, command preshaping techniques have only been applied to systems with actuators capable of producing a continuum of different command levels. In this context, particular attention has been given to the space shuttle Remote Manipulator System (RMS), an extremely flexible six degree-of-freedom robot arm used by astronauts as part of the shuttle's

Payload Deployment and Retrieval System (PDRS). This thesis will investigate the application of command preshaping techniques to systems where the actuators are only capable of producing two discrete command levels: on and off. The shuttle RMS will continue to be the focus of application, but this time the approach will be from a different angle: we will keep the arm still while we look at controlling the orbiter's attitude using its on-off Rotation Control System (RCS) jets without exciting undue vibration in the RMS or its payload.

What makes this problem interesting is the combination of the limitations posed by the constant amplitude command restriction and the wide spectrum of systems to which this could be applied. The use of "bang-bang" control means that some way must be found to handle command sequences with impulses of varying magnitudes. It also requires us to frame the problem in terms of a fixed amount of time for which the controller must be set to "on" (henceforth referred to as the "burn-time" or "on-time") in order to achieve the total desired impetus. Instead of convolving impulse sequences with arbitrary, continuous commands we must convert the impulses to pulses and divide them up in order to achieve the necessary burn-time without exciting vibration in the system. No longer do we have the ability to scale our commands to achieve the desired effect; our only degree of freedom is in their timing. At the same time, the number of real systems which must operate under these restrictions is enormous. In space, where lightweight, flexible structures are ubiquitous and external sources of damping are notoriously absent, nearly every system uses non-throttleable thrusters for propulsion, and the vast majority use them for attitude control as well. Command preshaping techniques would be a great boon to many of these systems, but current methods do not account for their

unique needs. Systems with constant amplitude actuators abound in terrestrial industry as well, and preshapers designed for them could no doubt be applied to good effect in a variety of disciplines.

The roots of command preshaping go back as far as 1958, to the posicast control system proposed by O.J.M. Smith. [26] The posicast system suffered from a lack of robustness to frequency error, however, and it was not until 1988 that the field truly got going with Singer's doctoral thesis. [24] Singer's application of constraint equations to solve for sequences producing no residual vibration and offering robustness to frequency error was what was needed to make the concept practical. As well as examining command preshaping in several different domains to broaden understanding of its principles, Singer provided solutions for handling systems with damping and multiple frequencies. Singer's work is summed up in [23], which also gives a good overview of other work that has been done in the controls field in terms of reducing vibration through the shaping of command inputs.

Singer's work has been continued by a variety of others. Singhose [25] described methods for generating more robust sequences by relaxing the zero vibration constraint at the design frequency. Hyde [11] expanded on the use of numerical optimization techniques for sequence design. Rappole [17] proposed the use of symmetric sequences for coping with multiple frequencies which offer substantial robustness with a smaller time-delay than sequences generated by convolution. Bong Wie apparently arrived at similar results to Singer's without realizing that work had already been done in the field. [31]

Researchers have applied command preshaping techniques to a variety of systems; a short list of accomplishments in the field can be found in [17].

Most of the previous work that has been done with command preshaping has relied on numerical optimization techniques for sequence selection. In this thesis a somewhat different approach will be taken. Instead of feeding a set of constraints into a nonlinear optimizer, we will develop a series of different base sequences and present a methodology for combining them to achieve the desired results. There are three major reasons for taking this approach to the problem. The first is one of simple realism: the flight computers on the space shuttle are extremely limited in power and memory, rendering sequence selection by numerical optimization impractical. The second reason is that constant amplitude preshaping is an area that is not yet well understood; by taking our solutions from a "black box" such as a nonlinear optimizer, we sacrifice the insight into the problem that can be gained by proceeding from basic principles. If we forgo the optimizer we can see *why* a particular sequence works well in addition to how well it works. This is closely bound up with the third reason, which is that until we understand the principles at work in constant amplitude preshaping, we don't really know the proper way to pose the problem to the optimizer. And if we don't give the optimizer the proper constraints, we won't get out the optimal results. Numerical optimization may be better suited to this problem once we have established the operating principles and can work out simple problems by hand to check the computer's results, but for the purposes of this thesis we have not yet reached that point.

Chapter 2 will provide a description of the shuttle-RMS system. It will begin with a discussion of the current system used for attitude control, the Digital AutoPilot (DAP). Next will come a brief description of the RMS. With the following section, which goes over the types of maneuvers which might be called for when a payload is extended on the arm, we begin to see the relevance of command preshaping techniques. This is then reinforced by an analysis of the types of problems which can result from interactions between the autopilot and RMS. A final section describes the solutions which have been developed so far for those problems, where they have succeeded and where they leave something to be desired.

Chapter 3 presents the theory of constant amplitude preshaping. It begins with a description of the simple mass-spring system used for the theoretical development and a discussion of the assumptions which are made. Proceeding step by step, we then look in turn at the response to a single pulse, two-pulse sequences, three-pulse sequences, and so on. Each section will look at issues of pulse width, insensitivity to frequency error, and time cost. New sequences specially designed for constant amplitude systems are introduced, as well as sequences capable of minimizing system loads. Issues involving more complex systems with damping and multiple frequencies of vibration are then addressed. A final section proposes a systematic method for selecting a sequence to match the needs of a particular application.

Chapter 4 takes the theory developed in the previous chapter and applies it to the shuttle-RMS system. It begins with a description of the software used to simulate the system and the particular configuration used to generate

experimental results. It then proceeds with a discussion of the methods used for frequency identification before going on to present the simulation results. Those results are presented for each of three payloads: the Upper Atmospheric Research Satellite (UARS), the Hubble Space Telescope (HST), and space station model MB6, as well as for an unloaded arm. Each payload is used to illustrate specific aspects of the theory presented in Chapter 3.

Chapter 5 concludes the thesis with an overview of the results and some suggestions for future work.

Shuttle Operations

Chapter 2

Later chapters in this thesis will focus on the theory and application of command preshaping. They will contain large numbers of graphs, equations, and diagrams, all of which will be presented with the ultimate aim of showing how one particular method of generating commands can be used to minimize vibration in a specific subclass of flexible systems. They will, I hope, do an excellent job of explaining the *hows* of constant amplitude preshaping. The purpose of this chapter is to explain the *whys*.

In particular, an argument will be made for why command preshaping techniques are needed in the space systems of today and in those of the future. It will do this by taking as an example the most complex space system currently in existence, the space shuttle. The chapter will begin by describing the digital autopilot (DAP) currently used to control the shuttle while in orbit. It will then proceed to discuss common maneuvers that the shuttle must be able to carry out while the RMS is in use. At this point a

description of some of the difficulties which may be encountered in these operations will point out the need for some sort of modification to the basic DAP. The current solutions to these problems will then be described, and the extent to which they succeed in their objectives will be assessed. The chapter will conclude by showing how command preshaping techniques could be used to address the problems encountered in present day shuttle-RMS operations and perhaps enable new tasks which could not safely be accomplished under the current system.

2.1 The Digital AutoPilot

The space shuttle on-orbit digital autopilot has evolved to its present state over a period of three decades. Though actual work on the shuttle's control system did not begin until the mid 1970's, the shuttle DAP took as its basis the system used by the Apollo project's Lunar Excursion Module, which began development at the MIT Instrumentation Lab in 1963. The initial version of the shuttle's flight control system was in development all the way through 1982, and was not completed until after five shuttle missions had flown. The long development time was in part due to the difficulty of compressing the entire package into the 128 kilobytes (K) of RAM available in the general purpose computer (GPC). The first studies analyzing dynamic interactions between the DAP and the RMS were not conducted until the mid-1980's, with work on the recent alt-mode and notch filter upgrades beginning only in 1989 after the GPC's memory capacity had been upgraded to 256 K.

The shuttle uses a system of non-throttleable thrusters to maintain attitude and conduct maneuvers on-orbit. The system consists of the original 38 primary jets (PRCS) and the 6 vernier jets (VRCS) which were added when it was realized that finer control than the primaries could provide would sometimes be needed. All of the jets may be switched on and off at the autopilot's cycle rate of 12.5 Hz, giving a minimum impulse firing duration of 80 milliseconds. The primaries, producing 870 lbs of thrust and originally intended for use in all situations, are still used for normal shuttle maneuvers but are often superseded by the verniers for delicate work. They are distributed around the orbiter in 14 groups, as shown in Figure 2.1, and produce angular accelerations of a bit under 1 deg/sec^2 with an unloaded orbiter. The 6 vernier jets, each producing 24 lbs of thrust, are used during RMS operations and whenever tight control is needed. They produce accelerations roughly 3% the size of their primary equivalents. Their main disadvantage is a lack of redundancy; as they were not a part of the original design they are few in number and located haphazardly. Only two of the six jets can fail without loss of control. Because there is no control mode which allows a mixture of vernier and primary jet firings, attitude control falls solely on the primaries if one of the four critical vernier jets should fail.

Maneuvers can be conducted under the DAP in several ways. [6] Manual control can be performed in either a pulse-based or continuous acceleration mode. A semiautomatic control mode is also available which will cause the orbiter to rotate about an axis at a preset discrete rate when the rotational hand controller is deflected. Finally, maneuvers may be conducted automatically using the DAP's 3-axis phase plane control system.

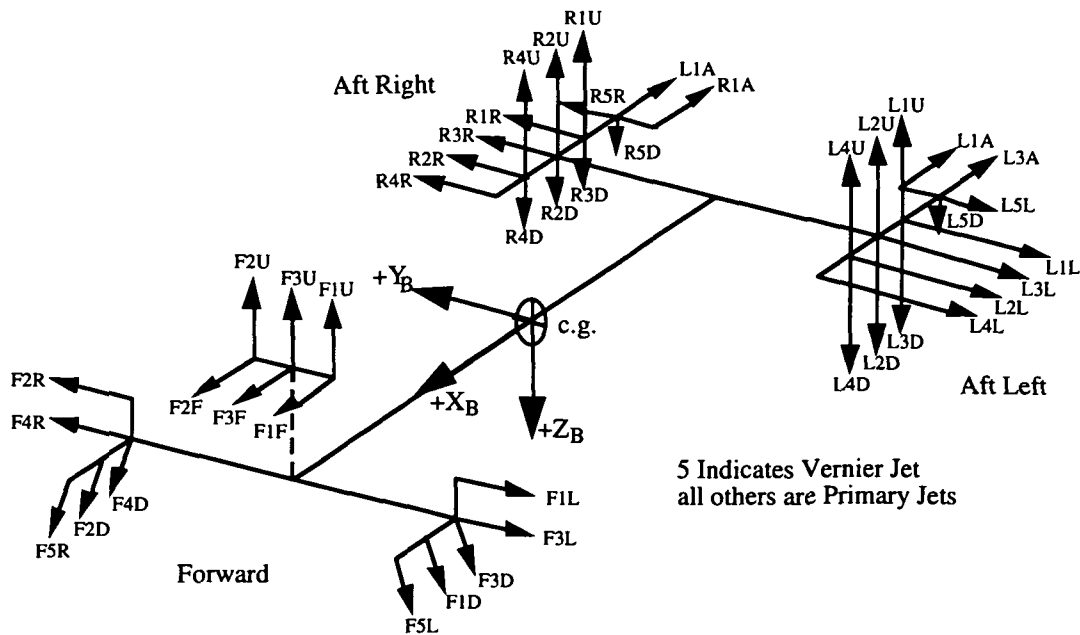


Figure 2.1: Locations and directions of the shuttle's primary and vernier jets.

When in automatic mode, the DAP tracks the orbiter's attitude on three separate phase planes, one for each rotational axis. For each axis a dead zone centered about the orbiter's desired position is defined and takes the shape shown in Figure 2.2. When the error moves out of the dead zone either by exceeding the rate limit or the attitude deadband, jets are fired to bring the attitude back into line. In normal operations this will ultimately result in a long period limit cycle about the desired attitude. The drift channels shown in the figure help to prevent excess firings and fuel waste. Maneuvers are conducted by precalculating a set of way points based on the specified maneuver rate. Each phase plane is then moved along these way points so that divergence from the desired trajectory will be corrected. When the orbiter enters the vicinity of the final desired attitude, the phase planes are centered on that attitude and normal station keeping resumes.

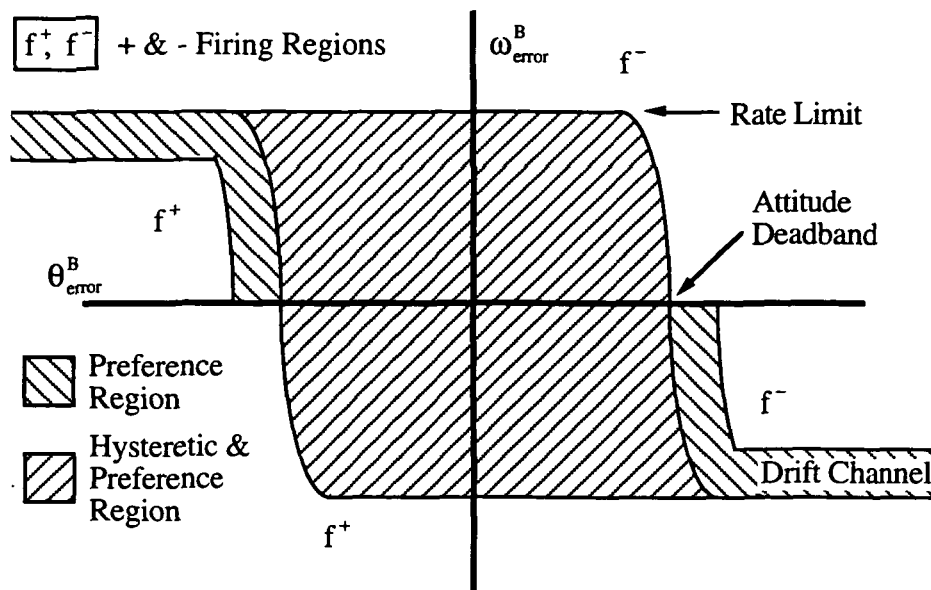


Figure 2.2: Phase plane deadbands for a single axis. Jet firings occur whenever the orbiter's position is outside of the shaded portion of the phase plane.

The behavior of the DAP is actually substantially more complex than the above description may imply. For more detailed information, see [7].

2.2 The Remote Manipulator System

The Remote Manipulator System is the robotic arm which serves as the major active component of the shuttle's Payload Deployment and Retrieval System (PDRS). Its main purpose is to deploy payloads from the orbiter's cargo bay and to capture payloads for retrieval or repair. It is also often used as a platform for astronauts conducting extravehicular activity. The arm has seven joints, one of which (the swingout) is normally fixed, giving it a total of six degrees of freedom. Its low weight (<1000 lbs) and great length (50 ft. 3 in.) combine to make it extremely flexible, especially when it is grappling heavy payloads.

The RMS is controlled from a panel in the aft section of the orbiter's flight deck. A variety of different control modes are available, but since our primary concern in this thesis is with the interaction between the arm and the shuttle autopilot, we will only consider the case where the joints are fixed in place and the brakes are on. For a discussion of the application of command preshaping to control of the RMS, see [24]. For more detail on the RMS and its control modes see [21].

2.3 Maneuvers

In order to determine how command preshaping might improve shuttle operations with the arm extended, we must first establish what kind of commands are typically given. In general, jet firings during RMS extended operations come from either planned maneuvers, DAP moding shifts, automatic station keeping, or are unanticipated. Each of these types of firings will be discussed in turn below.

Shuttle flight planners try to keep attitude maneuvers during RMS extended operations to a minimum. Sometimes such maneuvers are necessary, however; during STS-31, the mission in which the Hubble Space Telescope was initially deployed, seven separate maneuvers were planned while the arm was extended. [9] Some of these maneuvers require quite significant changes in attitude. For example, in STS-37 after the Gamma Ray Observatory (GRO) was deployed on the end of the arm it was necessary to change the orbiter's attitude so that the satellite's solar panels could be used to charge its batteries.

The orientations of the orbiter before and after this maneuver are shown in Figure 2.3. [28]

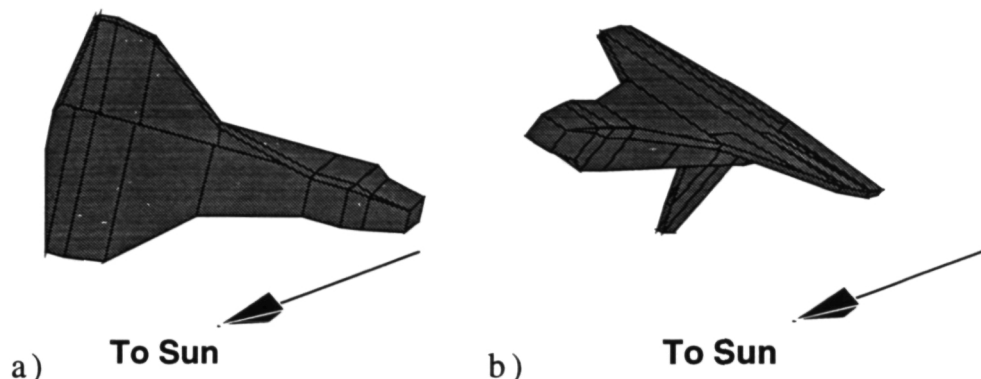


Figure 2.3: Orbiter attitude in two phases of deployment operations for the Gamma Ray Observatory. a) Deploy Attitude; b) Battery Charging Attitude

Maneuvers can also occur when the DAP is shifted from one mode to another. When RMS operations begin, the autopilot is generally put into free drift. As operations proceed, arm motion induces rotation in the orbiter, sometimes building up rates of .1 degrees/second or more. When RMS operations are halted and the autopilot is again engaged, attitude errors accumulated during the period of free drift must be corrected. These corrections are generally on the order of 10° — enough to require acceleration all the way to the maneuver rate before braking begins. This is the most common type of maneuver used while a payload is on the end of the arm.

Jet firings also occur with some regularity as a part of normal station keeping operations. Station keeping limit cycles ideally occur with a period of 10-12 minutes, but with a heavy payload out on the end of the arm cross-coupling between orbiter axes can result in much more frequent firings. In STS-31, 2-4 minute limit cycles were observed while the orbiter was in attitude hold with HST out on the end of the arm. [9]

Of course, not all maneuvers conducted during a mission can be predicted ahead of time. The orbiter's crew must be ready to react to a variety of contingencies, and confusion or miscommunication between crew and ground control personnel can exacerbate minor difficulties. In STS-31 three sizable unplanned maneuvers were conducted during HST deploy operations because of such miscommunication. [32]

2.4 Concerns with Dynamic Interactions

The reason that orbiter attitude maneuvers are of such concern is that significant problems can arise from the interaction between the dynamics of the orbiter and the various flexible objects connected to it or affected by its thruster plumes. Study of such dynamic interactions began as early as 1980, and flight testing to determine the actual extent of the problem was conducted as early as STS-8. Dynamic interactions can bring about undesirable behavior in the shuttle system in a variety of ways. Manifestations of these problems will tend to vary in magnitude depending on the payload mass, and can develop as short-period limit cycling, RMS brake slippage, or even catastrophic structural failure.

A number of different parts of the orbiter-payload system exhibit flexibility that can be excited by DAP attitude maneuvers. The one most often discussed is, of course, the RMS, but other components have problems with flexibility as well. The payloads on the end of the arm often have flexible components such as solar arrays or long antennae. The connections which keep payloads berthed to the cargo bay often have flexible components which

could cause problems. A payload deployment mechanism which has been used to roll out Intelsat's satellites has been the target of dynamic interaction studies. A slightly more indirect interaction which has caused much recent concern involves the space station's solar arrays. The plumes from the orbiter's jet firings could excite serious vibration in these large, fragile structures.

Every time one of the shuttle's RCS jets fires it induces some level of vibration in each of these locations. In addition to the difficulties associated with unwanted motion, vibration from jet firings can produce two other types of serious problems. One of these problems is the potential for excessive loading to fragile satellite components such as solar arrays. The other is the possibility of a resonant cycle developing if jet firings occur near an arm or payload frequency.

Unwanted motion can cause difficulties in a number of different ways. Vibration in the RMS makes it difficult to get good video clips from the cameras mounted on the arm. Astronauts using the RMS as a platform often find their work disrupted by vibration in the arm and then delayed further as they wait for the oscillations to die down. Vibration from jet firings can also disrupt the pointing of sensitive instruments or even cause a collision if RMS operations are being conducted in an environment with little clearance.

Loading internal to the RMS can (and often does) cause the brakes which hold the arm's joints in place to slip. Brake slippage can send the arm into a singularity, which is dangerous for any sort of robot. It can also cause a joint angle to reach its hard stop value, at which point further loads could damage

the arm. Finally, brake slippage can cause the arm or its payload to collide with another object, possibly damaging the orbiter or expensive equipment.

In addition to the dangers of collision, payloads are sometimes directly at risk from excessive loading. Satellite solar arrays are a common candidate for this sort of risk, as their large area and usually flimsy structure combine to make them quite fragile. Loads induced by shuttle jet firings could easily cause these structures to bend or break, ruining an expensive satellite. This danger is particularly high if one of the vernier jets should fail, requiring a switch to PRCS control.

There is one more important manifestation of the dynamic interaction between the autopilot and the flexible structures attached to the shuttle. When vibration occurs at certain frequencies it can resonate with the shuttle's automatic control system in such a way as to render it unstable. This may happen when a payload has an internal frequency or is heavy enough to alter one of the arm's frequencies such that it falls near the break-point of the DAP's state estimator. When this occurs the combination of phase shift and magnitude attenuation in the estimator can result in a rapid series of jet firings which send the attitude from one rate limit to the other again and again in what is known as a short-period limit cycle. [2] This behavior can amplify the initial vibration, increasing loads to the point where the brakes begin to slip in the RMS and damage to the arm or the payload becomes possible. Such behavior has never been observed in an actual flight but has been demonstrated repeatedly in high-fidelity simulation with payloads such as the Gamma Ray Observatory (GRO) and the Hubble Space Telescope (HST).

The danger posed by these load issues and possible instabilities varies sharply with the mass of the payload. An unloaded arm is unlikely to experience any difficulties except for possible camera jitter. With smaller payloads instabilities may not produce high enough loads to cause brake slippage, but could result in substantial fuel waste from the large number of firings. They may also have flexible appendages which could be vulnerable to the loads from primary jet firings. Midsize to large payloads such as GRO and HST may be troublemakers, especially if they have flexible components. The nominal limit for payload size on the RMS is 65,000 lbs. This may be exceeded during construction and maintenance of the space station, in which case great care will be needed to avoid excessive vibration and damage to the arm.

2.5 Current solutions

The dangers outlined in the section above have resulted in a situation where a large amount of pre-flight analysis is conducted for each mission where dynamic interactions could be of concern. The analysis seeks to ensure that even in a worst case situation the mission's success will not be threatened. Testing is done for jets firing in the worst possible direction and analysis is conducted for PRCS control as well as VRCS in case one of the vernier jets should fail. Over the years since the shuttle's initial flight, this analysis has indicated that for some payloads the original flight control system is inadequate to ensure mission safety. As a result, in 1989 Draper Laboratory engineers began work on a series of upgrades designed to take advantage of the GPC upgrade (to 256 K of RAM) due in 1991. The upgrades would seek to

remove system instabilities which might result in short period limit cycling and keep loads down in the event of a VRCS failure.

The first of these upgrades, known as the alt-mode, first flew in 1992. The alt-mode is intended to provide a way of firing the orbiter's primary jets without inducing the large loads which could damage a payload's flexible components. It does this by requiring that jet firings be of minimum duration and separated by some minimum amount of time. This separation would ideally be set to one half of the main system period, but for reasons involving the shape of the attitude and rate deadbands it is often set to 3/4 of a period instead. Under ideal conditions, maneuvers conducted under alt-mode produce about the same vibration and loads levels as maneuvers conducted using the VRCS.

Unfortunately, the alt-mode is by no means an ideal solution to the problem it addresses. The combination of minimum impulse firings and large delays gives very poor control authority. In partial compensation for this reduced authority, the alt-mode should achieve the same effect as the simplest impulse preshaping sequences by firing half-period intervals. However, the alt-mode is integrated into the phase plane controller system in such a way that the specified delay is only the *minimum* spacing between firings. Firings can occur at larger time intervals, and could conceivably occur in such a way as to aggravate the problem instead of fixing it, such as firing at full period intervals. There is no way of knowing ahead of time. The alt-mode also suffers from an inability to avoid exciting more than a single frequency. Delay times have to be selected through a painstaking process of trial and error to make sure the firings do not aggravate any of the frequencies in the shuttle-RMS-

payload system. Finally, the large delays between jet firings cause even simple maneuvers to take a substantially longer time than they would under VRCS control. These delays can be costly in a shuttle mission where every minute is precious.

The second upgrade, due to fly for the first time in the summer of 1994, uses notch filters in the DAP's rate estimator to prevent instabilities from arising. [2] By removing from the rate estimate the frequencies at which flexible components might oscillate, the notch filters prevent the autopilot from getting into a pattern of jet firings which could result in short period limit cycling. Like the alt-mode, the notch filters require extensive pre-flight analysis to be of use.

2.6 The Role of Command Preshaping

The previous sections have provided an overview of the system in question. We have seen something of the limitations imposed on shuttle operations by dynamic interactions with its payloads and deployment mechanisms. We have also discussed some of the methods which are now used to get around those limitations. We can now look at the roles which could be played by a command preshaping system.

As we shall see in the next chapter, constant amplitude preshapers are capable of reducing a flexible system's vibration with high tolerance for errors in frequency estimation. Moreover, they are not limited to a single mode; preshapers can be designed to reduce vibration at a variety of frequencies with an appropriate level of robustness at each. They are also

capable of producing firing sequences which minimize the loading in a single-frequency system in a time-efficient manner. These features could be used to improve current shuttle operations and make possible new missions which could not be carried out with the original autopilot.

Impulse Preshaping in Constant Amplitude Systems

Chapter 3

In this chapter we will present the theory behind the application of impulse preshaping to constant amplitude systems. We will describe the difficulties which arise from the constant amplitude restriction and methods of overcoming those difficulties. We will show how the sequences and methods used in continuous systems can be applied to constant amplitude systems and then go on to add new sequences and methodologies to the existing library of preshaping techniques. In the process of examining impulse preshaping in constant amplitude systems we will generate some new insights into impulse preshaping in general.

The chapter begins with a description of the model upon which the subsequent analysis is based. It then proceeds to discuss that model's response to a single pulse, and how it differs from the response to an impulse. From the discussion of the behavior of a single pulse we continue onward to the

combination of two pulses to form our first simple preshaping sequence. This sequence is compared to the equivalent sequence in the continuous domain, and the changes in thinking necessitated by the constant amplitude framework are discussed. Having seen how the simplest impulse sequence is adapted to a constant amplitude system, we then take a look at a slightly more complex sequence. Here we again address the issue of the approximation of impulses with pulses, and how that approximation is affected by pulse width. Also in this section we discuss the difficulties of measuring performance in constant amplitude systems.

Having laid out the ways in which standard impulse sequences can be adapted to constant amplitude control systems, we then go on to describe some new sequences which are uniquely suited to constant amplitude applications but which may also prove useful in the continuous domain. We begin by describing sequences based on impulses placed at $1/3$ and $1/5$ period intervals instead of the normal half period. Next we show how judicious placement of pulses can be used to limit the loads induced by a command while keeping residual vibration low.

At this point we pause to take a closer look at the assumptions which the preceding analysis makes use of. The effects of damping and time discretization are described, and strategies for designing multiple-mode sequences are proposed. Finally, the chapter concludes with a discussion of how all these things can be integrated into a single sequence design process.

3.1 Theoretical Model

Our discussion of preshaping in this chapter will revolve around a very simple model. The flexible spacecraft to be controlled will be modeled as two equal masses connected by a massless spring. This configuration is shown in Figure 3.1. The force applied to one end is directed to the right and may have a magnitude of either zero or F . Two modes will be apparent in the motion of this system: a rigid body acceleration and a single flexible mode with frequency $\omega = \sqrt{k(M+m)/Mm}$, where k is the spring constant and M and m are the two masses.

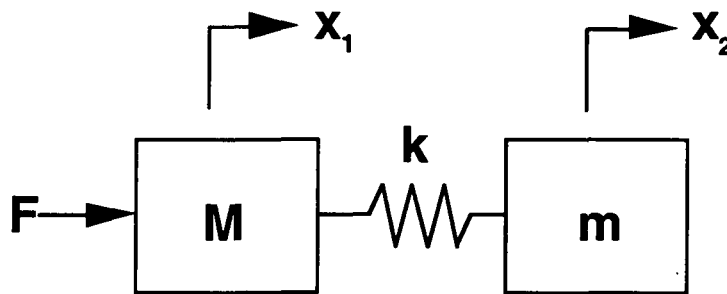


Figure 3.1: Two-mass, one spring model used for analysis in this chapter. The system is acted on by an external force F and has a spring constant k .

Several assumptions and simplifications are inherent in this model of a flexible spacecraft. The model is limited to a single flexible mode, damping is not included, and any number of nonlinear effects likely to be present in a real space system are ignored. All of these issues will be addressed in later sections.

The derivations and equations to follow also assume that we can command pulses to begin at precise intervals; we may want a pair of pulses separated by precisely one half period or seek small adjustments to one side or another.

Real systems operate on digital command cycles of finite length, however. This means that commands will almost always be implemented slightly before or after the ideal time.

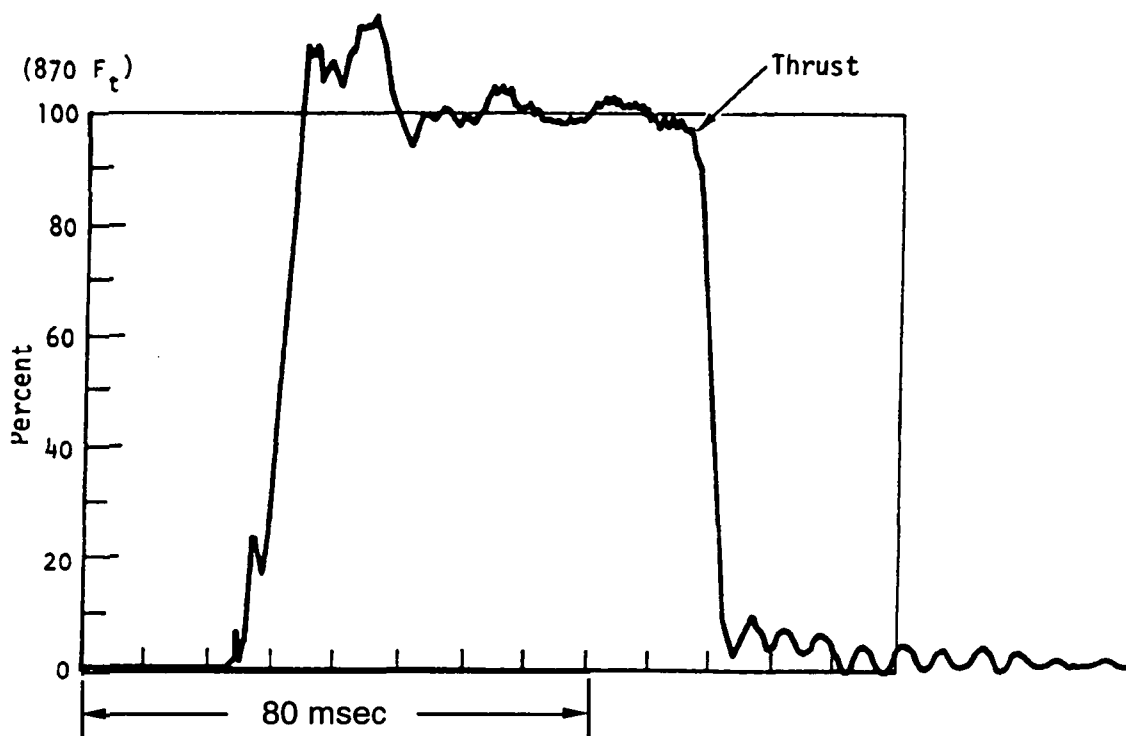


Figure 3.2: Typical thrust profile for a space shuttle Primary Rotation Control System (PRCS) jet firing a minimum duration (80 millisecond) pulse.

It is also worth pointing out that in a real system pulses do not come in perfect square waves. Figure 3.2 shows a “typical” thrust profile from one of the space shuttle’s PRCS jets. [22] This 80 millisecond pulse doesn’t reach its full thrust level for 30 milliseconds after the activation time and has an overshoot worth a good 3% or so of the total thrust provided by the pulse. The relative effect of this overshoot will vary with the pulse size, and ambient conditions may affect the pulse profile.

3.2 The Single Pulse

Taking in hand, but for the present ignoring all the qualifiers described above, the equations of motion for the system described above are well known. The system's time response (in the flexible mode) to a pulse beginning at t_1 and ending at t_2 will be of the form

$$\chi(t) = A(\cos \omega(t - t_1) - \cos \omega(t - t_2)) \quad (3.1)$$

where A is a constant dependent upon the system parameters and shall henceforth be taken to be unity.

This simple force profile is shown in a variety of forms in Figure 3.3. In (a) we see an impulse based representation of the command, where the entire change in velocity is imparted in a single instant. What is shown in (b) is a pulse of finite width $\Delta = t_2 - t_1$ centered about $t = 0$. If we shift the time scale so that the pulse in (b) begins at $t = 0$, then we can say it approximates an impulse of equal area placed at $t = \Delta$. This approximation improves in the limit as Δ/T goes to zero. The picture in (c) is what shall be called hereafter a pulse sector diagram, as opposed to the pulse vector diagrams presented in Singer [24] and Singhose [25]. It shows the duration of the pulse in relation to the period T (360°) of the system's flexible mode.

It is the last of the pictures shown in Figure 3.3 that is the most important to remember, however. The graph shown in (d) comes from an analysis of how the residual vibration changes as a function of the pulse width. Analytically, we can find this function by expressing Equation 3.1 as

$$\chi(t) = B \cos(\omega t + \phi) \quad (3.2)$$

and then solving for the residual amplitude B , giving us the expression

$$B = \sqrt{2(1 - \cos t_p)} \quad (3.3)$$

where t_p is the pulse width. Looking at the curve, we can see that the maximum residual amplitude will occur for a pulse width of $T/2$ and that for a pulse width of T there will be no residual at all.

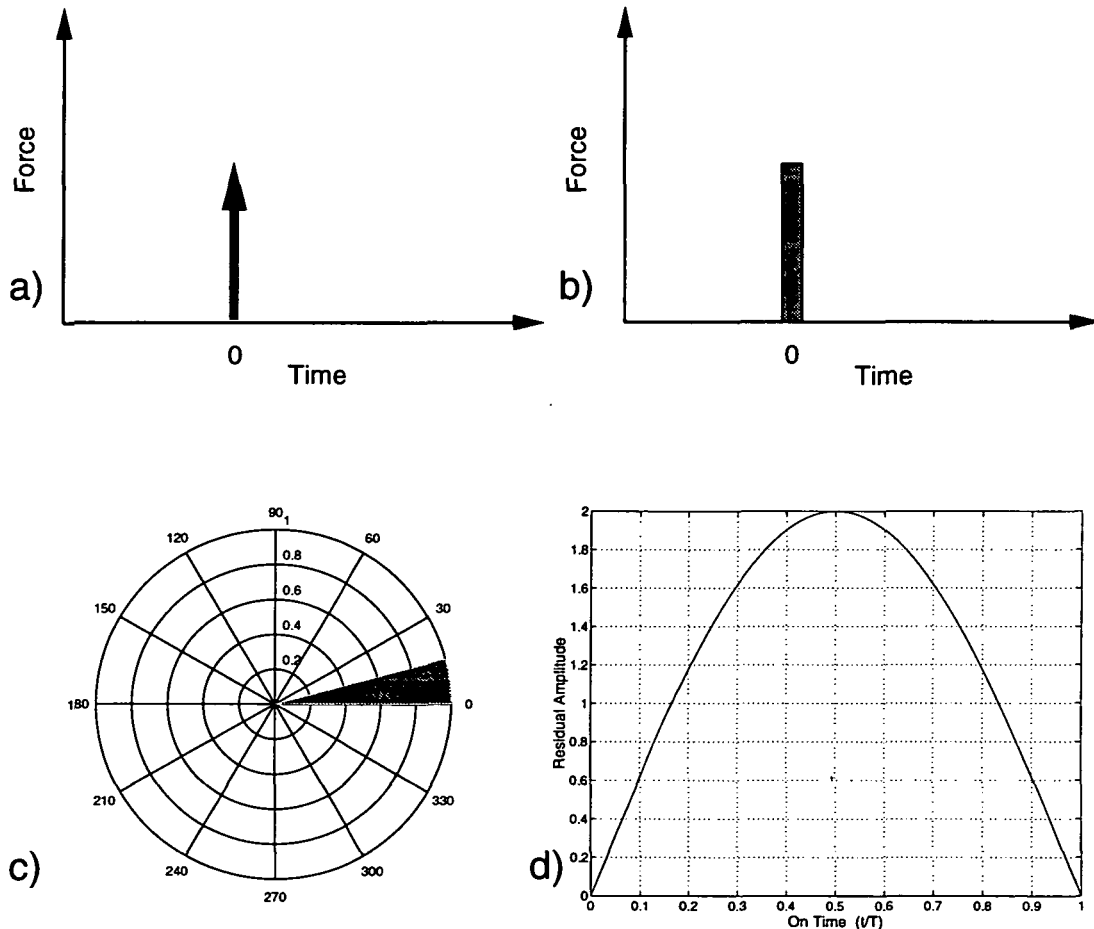


Figure 3.3: The single pulse. a) Impulse representation (continuous domain); b) Pulse representation (constant amplitude domain); c) Pulse sector diagram; d) Residual vibration as a function of command size

This dependence of residual amplitude on pulse width makes measurement of the performance of preshaping sequences more difficult. We would like to be able to measure the residual amplitude from a preshaping sequence as a percentage of the residual left from a single pulse of equal total width. If we do this, however, we find that our performance measure is dependent on the command size. When the total on time (t_b) approaches T , even an extremely small residual amplitude left by a preshaping sequence may seem large compared to the residual left by a single pulse.

The desire for an unbiased measure of a sequence's performance forces us to use a stick without a direct connection to the system in question. Instead of normalizing by the residual amplitude left by a single pulse of equal total magnitude we must divide by the residual amplitude left by an *impulse* of equal total magnitude. In one sense, this comparison is devoid of physical meaning, as such an equivalent impulse cannot exist in a constant amplitude system.

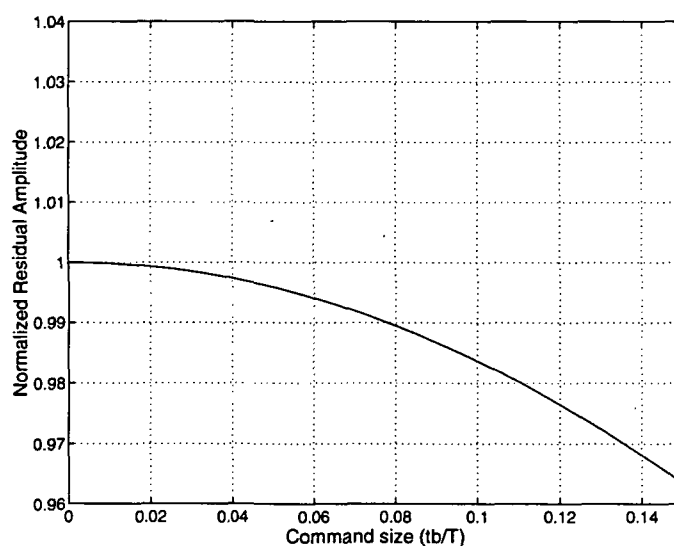


Figure 3.4: Graph of the residual amplitude from a single pulse normalized by the command size vs. the command size. Shows that the residual scales linearly when the command size is small compared to the system's period.

If we look at it in another way, however, the meaning is still there. For if the command size is small relative to the system's period T , we can measure the residual of a sequence as a fraction of the residual left by an equal width single *pulse*. The reason this works can be seen by looking at the graph in Figure 3.4. It shows the residual amplitude from a single pulse of varying width, just as in Figure 3.3 (d), only this time normalized by the pulse width. For small pulses, this normalized residual can be approximated as constant at 1. Thus the residual from a single, small pulse will be twice as large as the residual from a single pulse of half the duration, something which can not be said of a larger pulse. In this regime, the amplitude of the response to a pulse is equivalent to that of an impulse of equal magnitude.

3.3 The Line Sequence

As has been shown in Singer, the simplest way to achieve zero residual vibration is to use a two pulse sequence. [24] In a system with no damping, this sequence is just two impulses of equal size spaced $T/2$ seconds apart. Figure 3.5 (a) shows the impulse diagram for the sequence, while (b) shows the sequence converted into finite length pulses. In (c) we have the sequence's pulse sector diagram, and in (d) the insensitivity curve.

It is here that we get our first look at using a preshaping sequence in a constant amplitude system. In this case, as both of the pulses are the same amplitude anyway, the constant amplitude restriction has no real adverse effect on the system. Changing the pulse width has no real effect either, as each incremental addition to the duration of the first pulse is exactly canceled

by an equal addition to the second. This is most easily understood by looking at the pulse sector diagram and noting that each radial vector within the first sector is exactly in line with, and canceled by, a radial vector in the second sector. This linear cancellation is why we will call the sequence a *line* (or *L*) sequence.

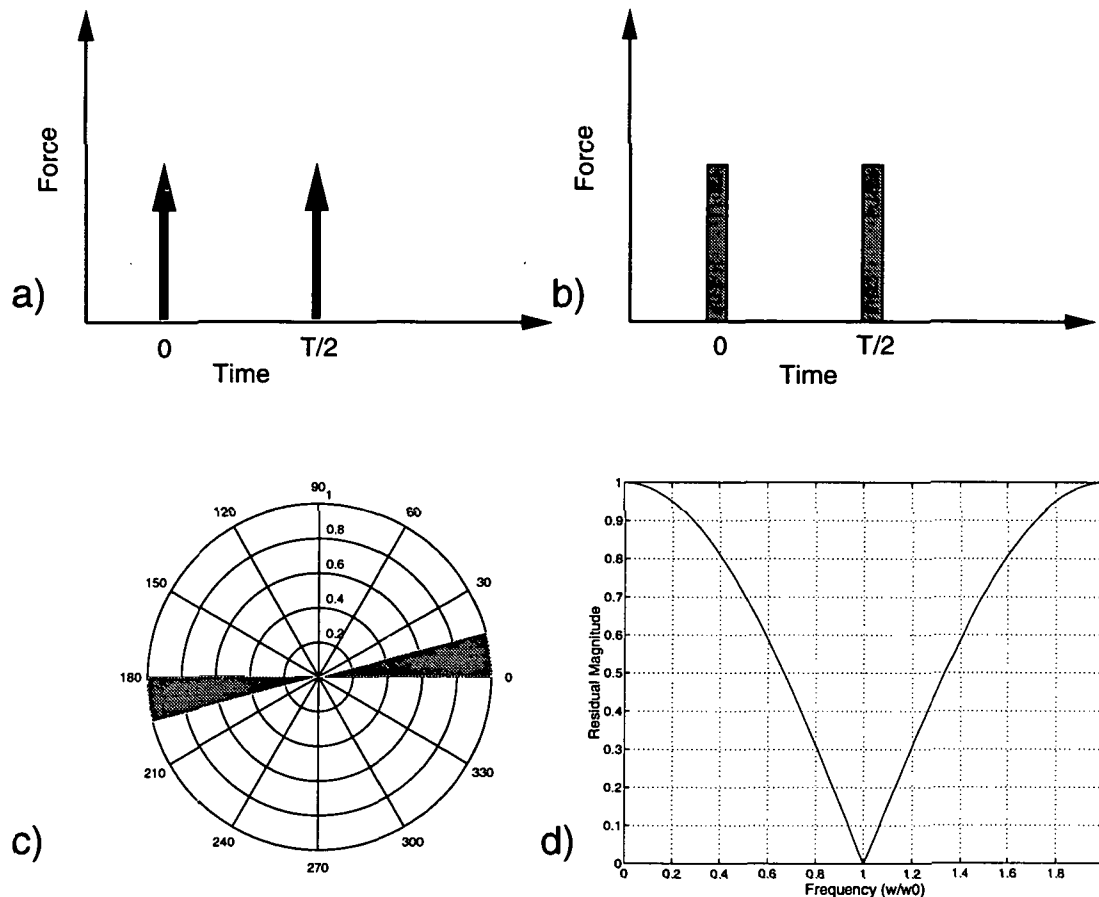


Figure 3.5: The line sequence. a) Impulse representation (continuous domain); b) Pulse representation (constant amplitude domain); c) Pulse sector diagram; d) Insensitivity curve

The insensitivity curve shown in Figure 3.5 (d) deserves some further comment. It was generated, as with all the insensitivity curves in this thesis, by approximating each pulse as an impulse and using vector summation to find the magnitude of the residual vibration. Since the impulse approximation

is equivalent to saying the command size is small, this is the same process as was described in the section above on the single pulse response. All of which means that we have a curve which is exactly identical to the insensitivity curve presented for this sequence by Singer but with an interpretation attached to it that makes more sense in the constant amplitude setting. We are in effect measuring the residual vibration produced by the sequence as a fraction of the amplitude of the vibration produced by an impulse of equal magnitude.

Having established what we mean by insensitivity, we can go one step further and derive an analytic formula for the insensitivity curve. Using simple trigonometric identities it can be shown that the amplitude of the response to two impulses at $t_1 = 0$ and $t_2 = \pi/\omega_d$ as a function of the actual system natural frequency ω_N is

$$r = \left| \cos \left(\frac{\pi \omega_N}{2 \omega_d} \right) \right|. \quad (3.4)$$

Analytic formulae for the residual amplitudes of other sequences will be presented in their respective sections.

3.3.1 Time Cost

Though we have just claimed that pulse width has no real effect on a line sequence, this is only true up to a point. As the pulse width increases, at some point the first pulse will run into the second, forming a single large pulse. This special case corresponds to the full period pulse discussed in Section 3.2. To increase the command size beyond this point we need to restart the two-

pulse sequence at $t = T$ as shown in Figure 3.6. This procedure can be repeated again and again, as necessary. This cyclic discontinuity doesn't really change the behavior of the sequence itself, however. What it does change is the way that time delay should be measured for constant amplitude sequences.

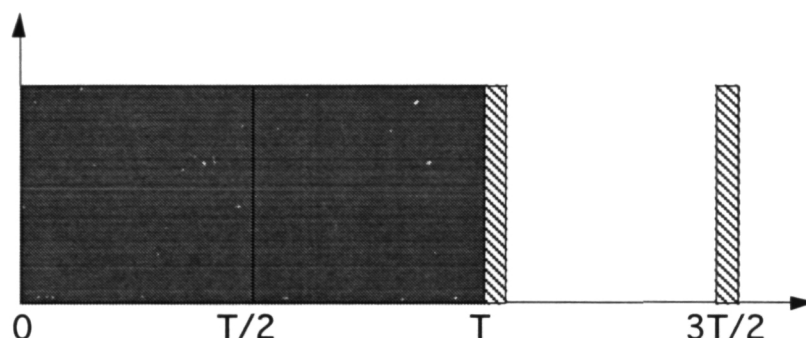


Figure 3.6: Force Profile for a line sequence requiring a burn time longer than one full period.

In a variable amplitude system preshaping is implemented by convolving impulse sequences with a command signal. Thus the preshaped command will always finish after the unshaped command by the length of the impulse sequence, allowing us to safely specify a single number as the time delay for a given impulse sequence. For the L-sequence, the time delay is $T/2$.

In a constant amplitude system, however, there is no fixed time delay. What we have instead is a time *cost*, a number measuring the additional time necessary to carry out a command using a preshaping sequence instead of a single pulse. If we measure it in this way, this time cost will be a function of the command size, and will look like the curve in Figure 3.7. It will start out equal to the variable amplitude time delay at $T/2$, but decrease linearly until it hits zero, at which point it will jump back up to $T/2$. As we shall see in later sections, more complex sequences exhibit the same type of behavior except

that the time cost does not reach zero, instead accumulating a fixed increment with each cycle.

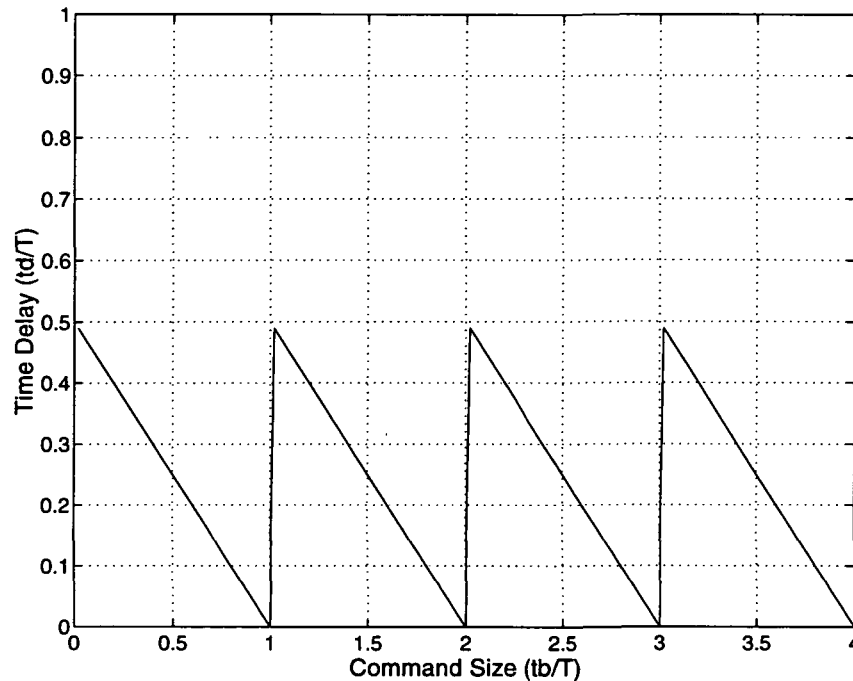


Figure 3.7: Time cost as a function of command size. No time penalty is incurred when the burn time is an integer multiple of the system's period, but the delay jumps up to its maximum value if the required burn time is a little bit longer.

3.4 The L^2 Sequence

The line sequence certainly has its place in command preshaping, but is limited by its lack of robustness to frequency error. As has been shown by Singer, the desired robustness can be added by convolving a line sequence with itself to give a three impulse sequence such as the one shown in Figure 3.8 (a). [24] To distinguish this sequence from other sequences with three impulses, we shall call it an L^2 sequence. The L^2 sequence addresses the lack of robustness of the line sequence, but at the same time its varying amplitudes create a host of other difficulties in a constant amplitude system.

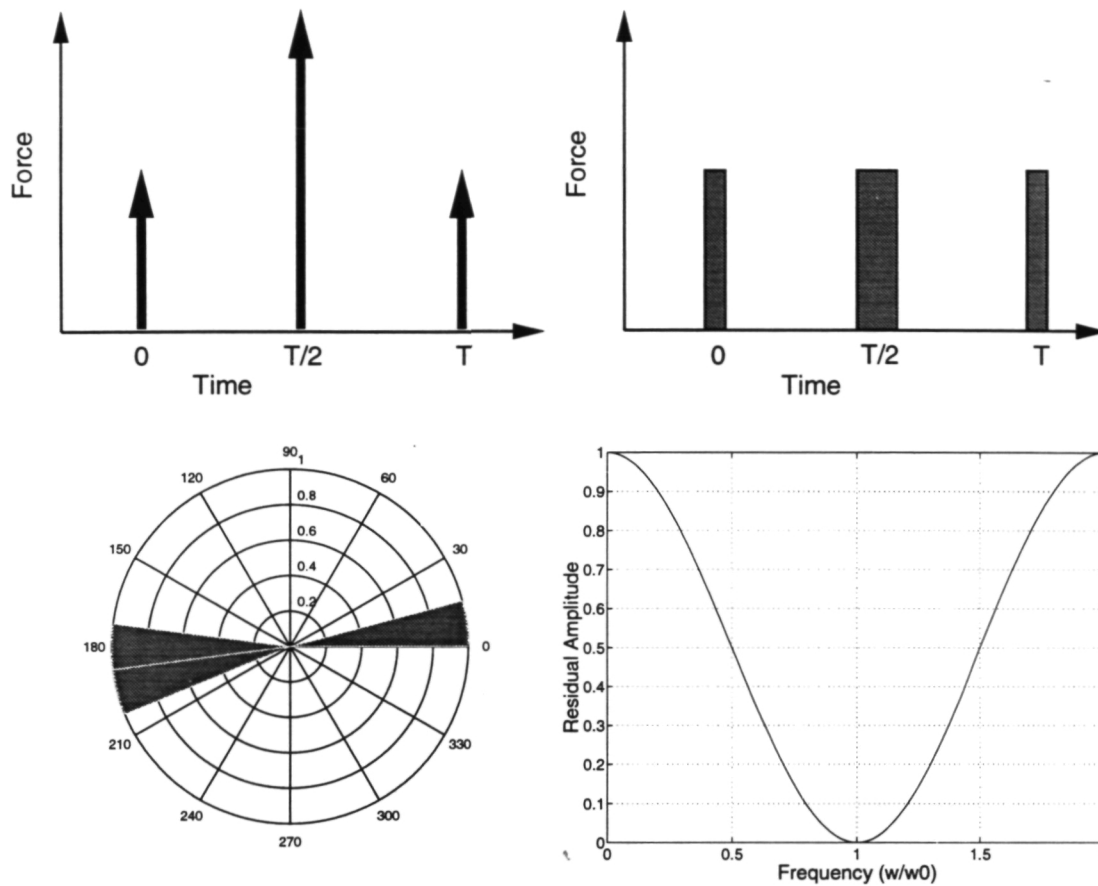


Figure 3.8: The L^2 sequence. a) Impulse representation (continuous domain); b) Pulse representation (constant amplitude domain); c) Pulse sector diagram; d) Insensitivity curve

The essential problem here is what to do with the center pulse. An obvious first guess is to approximate the double-height pulse as a single-height, double-width pulse centered on the same spot as the double height would be. This is shown in Figure 3.8 (b) as a time sequence, and in (c) as a pulse sector diagram. In (c) the fourth pulse overlaps the first.

Figure 3.8 (d) shows the insensitivity curve for the L^2 sequence. Using the same methods as above, we can show that the analytic formula describing this curve is

$$r = \cos^2 \left(\frac{\pi \omega_N}{2 \omega_d} \right). \quad (3.5)$$

At this point we need to stop for a moment and establish some terminology to avoid later confusion. In the pages to follow, there will be a lot of discussion of pulse sequences that contain a number of pulses of differing duration. The *width* of one of these sequences is the duration of the first pulse in the sequence, and will be represented at times by the symbol Δ . The individual pulses within the sequences have widths defined in terms of that initial pulse; a pulse of width 1 is equal in duration to the initial pulse, while a pulse of width 2 lasts for twice as long as the initial pulse. The *length* of the sequence is the time between the first and last pulses in the sequence. An L sequence has a length of $T/2$, while an L^2 sequence has a length of T .

So when we say we will approximate the center pulse of the L^2 sequence as a pulse of width 2 centered on the same spot as the double height pulse would be, we mean that the pulse will last twice as long as the first pulse and begin at $t = T/2 - \Delta/2$. The shift of $-\Delta/2$ comes because the first pulse, beginning at $t = 0$, is actually centered about $t = \Delta/2$. To center the middle pulse about $t = T/2 + \Delta/2$, we have to shift it a bit to the left.

3.4.1 Concatenation

Unfortunately, this simple approximation is only good when the width of the sequence is small compared to the period T . As the command size increases the sequence's width also increases, the approximation breaks down, and performance degrades. This may be avoided by concatenating several L^2 pulse sequences of smaller width to generate a larger command. Figure 3.9 shows how this process works. In (a) we see a simple L^2 sequence. In (b), another sequence of equal width has been placed next to the first. An off period of one pulse width is necessary between the first two pulses of the new sequence in order to keep the middle pulses from interfering with each other. When this process is used for larger commands, the force profile might look something like (c). Notice that the central pulses have all merged to form a single, large pulse. On either side the alternating sequence of on and off commands generates the same effect as a single pulse of half the amplitude. Using this technique, we can generate arbitrarily large commands that have the same robustness characteristics as impulse sequences in a continuous system. The penalty we pay for this performance is in the large number of switches required by the sequence, which can greatly add to the wear and tear on the command actuators. This added wear encourages us to investigate the effect of using wider pulses on preshaper performance.

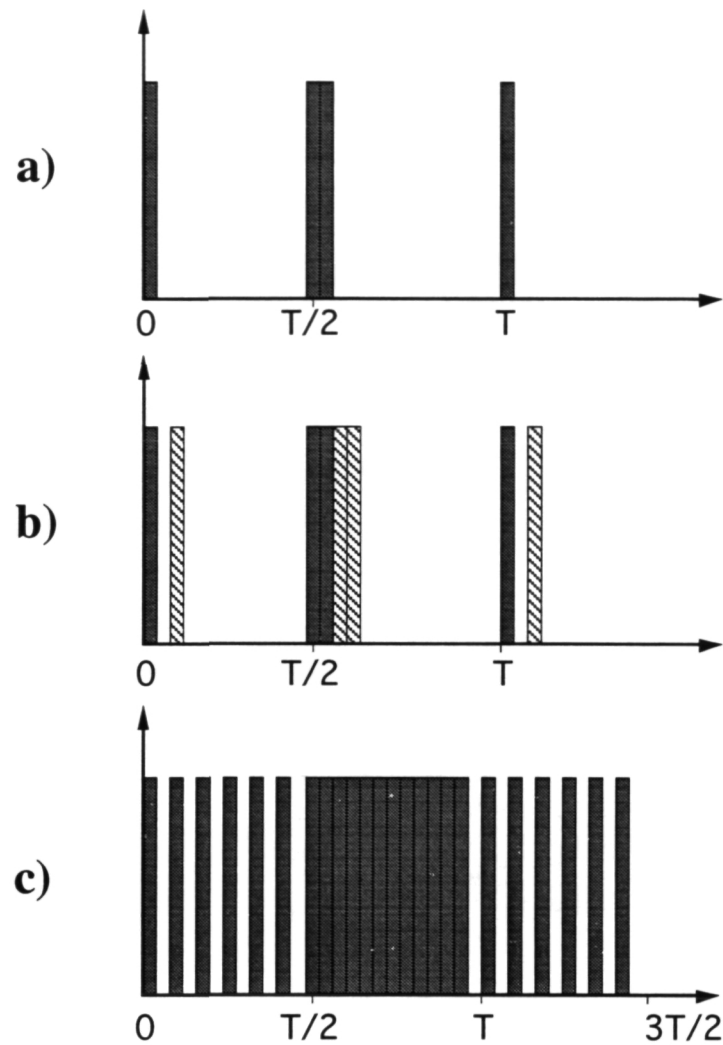


Figure 3.9: Construction of large sequences through concatenation. a) A basic constant amplitude approximation for an L^2 sequence; b) An additional L^2 sequence concatenated onto the first to provide an equal accuracy in approximation; c) The concatenation process taken to its logical extreme.

3.4.2 Pulse Width

As the pulse width increases, the residual vibration for an exactly known natural frequency gets larger and larger. Using a little trigonometry, we can find an exact formula for this dependence. We start by superimposing the responses from the individual pulses in the sequence to get

$$\chi_{L^2}(t) = 2 \cos \omega t - 2 \cos \omega(t - \Delta) - \cos \omega\left(t + \frac{\Delta}{2}\right) + \cos \omega\left(t - \frac{3\Delta}{2}\right),$$

from which, using a bit of trigonometric manipulation, we can extract the expression

$$(10 - 8 \cos \omega \frac{\Delta}{2} - 8 \cos \omega \Delta + 8 \cos \omega \frac{3\Delta}{2} - 2 \cos 2\omega \Delta)^{\frac{1}{2}} \cos(\omega t + \phi), \quad (3.6)$$

where ϕ is a phase shift that is unimportant for our purposes. All we are concerned with is the residual amplitude of the above expression, shown in Figure 3.10 normalized in two different fashions. The solid line shows the residual normalized by the residual amplitude from a single pulse of equal total width. The dashed line shows the residual amplitude normalized by the response to an impulse of equal total magnitude — the scheme which we use to generate the insensitivity curves measuring a sequence's effectiveness. The two normalizations produce equivalent results when the pulse width is small compared to the period T . But as the pulse width gets larger, the pulse response does not increase as quickly as the impulse response, and the two curves begin to diverge. When the sequence's pulse width reaches 12.5% of the period, the equivalent single pulse is a half period long and has reached its maximum response. After this juncture, the two curves begin to diverge more rapidly. In this regime the impulse normalization has effectively lost all meaning, while the equivalent pulse normalization approaches a singularity at a width of $T/4$. Forced to choose among poor alternatives, we find that the only scale which makes sense for measuring residuals from wide-pulse sequences is relative to the maximum response which can be generated from a single pulse — the peak of the graph in Figure 3.3 (d).

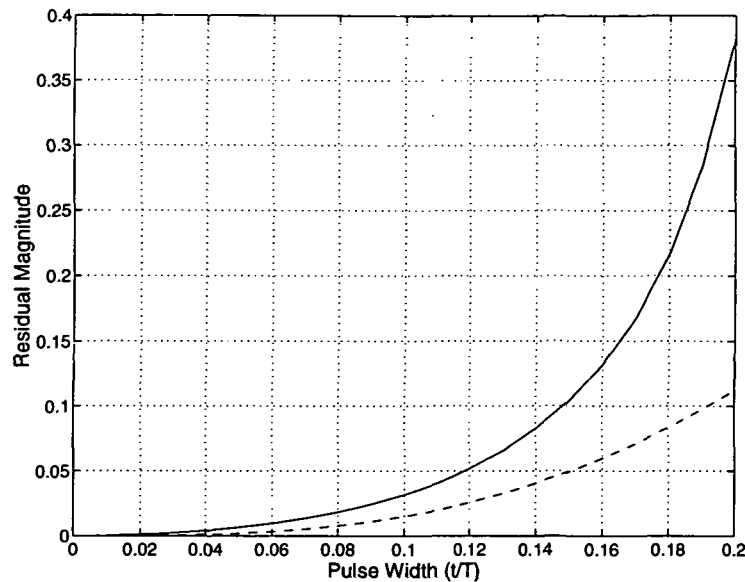


Figure 3.10: Residual magnitude at the design frequency for an L^2 sequence with varying pulse width. The dotted line uses an impulse normalization and the solid line uses a pulse normalization.

An understanding of these two different normalizations is key to the interpretation of the graphs shown in Figure 3.11. The figure presents various views of the effect of pulse width on the L^2 sequence's robustness characteristics. In the left column, the impulse normalization is used, while in the right column the residual amplitude is normalized by the maximum pulse response. In the top row we see a perspective view of the surfaces, in the middle a view from above with height differentials shown with shading, and in the bottom row we see contour plots which clearly delineate the shapes of the surface's features. In the contour plots the 5% mark is shown with a solid line, 20% with a dashed line, and 50% with a dash-dotted line. In all the graphs, the Y-axis (pulse width) extends only to a value of $T/3$. Greater pulse widths would cause the first and second pulses to run into each other.

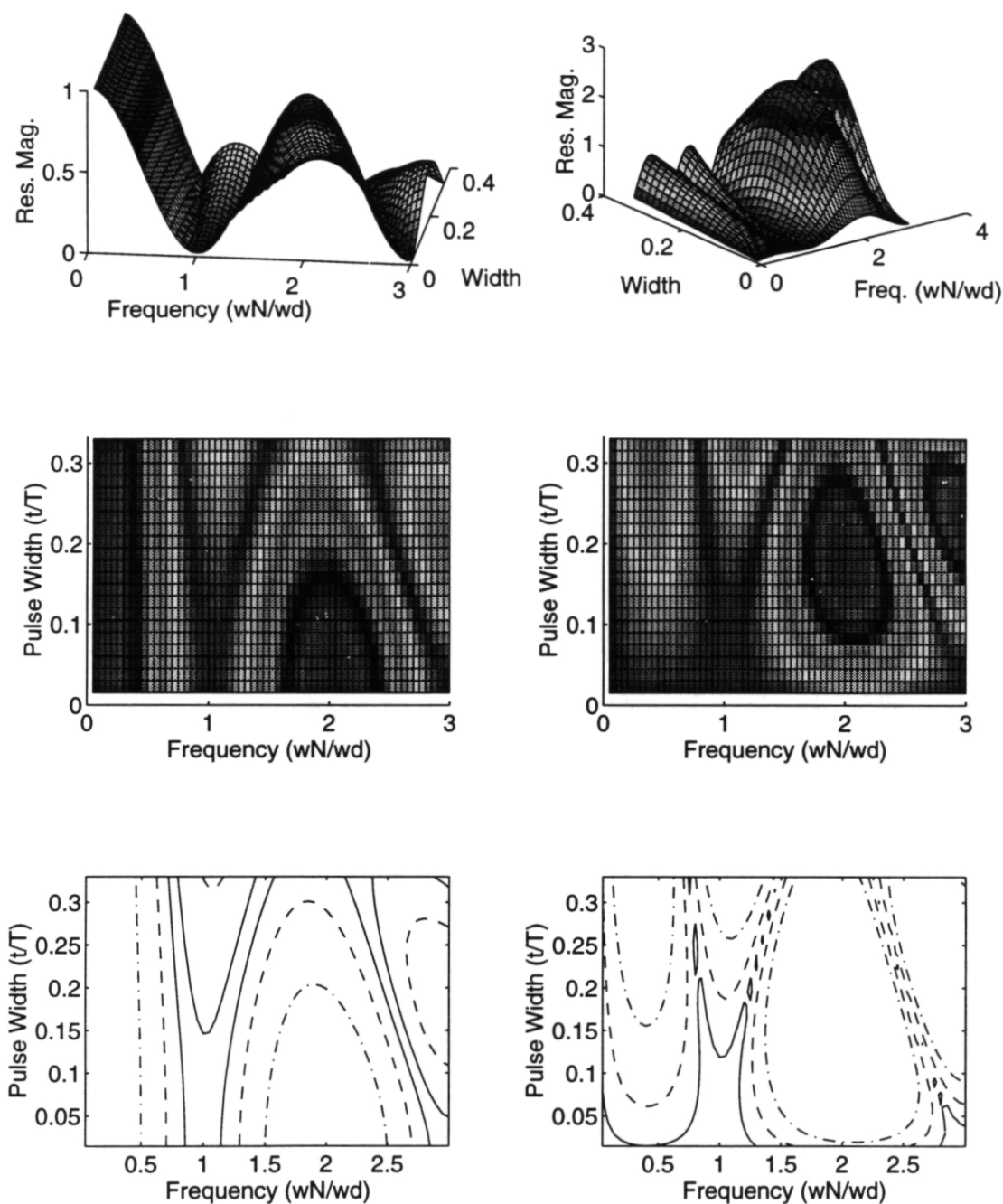


Figure 3.11: Residual magnitude for an L^2 sequence as a function of frequency and pulse width. Three different views are shown for each of two different normalizations. The graphs on the left use an impulse normalization while the graphs on the right use a pulse normalization.

One feature is common to both columns of Figure 3.11. As the pulse width increases the wide central valley at $\omega_N/\omega_d = 1$ splits into two narrower branches, visible as a dark 'V' in the central graphs. Since the split occurs in both columns, we know that this is not just an artifact of the normalization scheme used; at the floor of the valley the residual is near zero in absolute magnitude. This splitting phenomenon may be understood if one thinks of it in terms of pulse sector diagrams. For the case $\omega_N/\omega_d = 1$, the pulse sector diagram will look like Figure 3.8 (c). As the frequency changes, the first pulse will remain static on the diagram, but the third pulse, which overlaps with the first for $\omega_N/\omega_d = 1$, will begin to move. Eventually, the frequency may change enough (in either direction) so that the first and third pulses appear side by side in the pulse sector diagram and together cancel the second pulse exactly. The greater the width of the first pulse, the further the third pulse must shift to avoid overlap; this is why the two branches diverge from the centerline.

The splitting effect described above might lead one to think that the problem of adapting standard preshaping techniques to a constant amplitude framework has been solved. All one needs to do is adjust the design frequency a bit to take into account the pulse width and the residual is removed. This is a fallacy, however; regions of zero residual do exist for large pulse widths, but they are not nearly as robust as the one which approximates the standard L^2 sequence for small pulse widths. This may easily be seen by comparing the width of the central trunk in the contour plot with the two thinner regions which branch off at larger pulse widths. Clearly we can obtain the greatest robustness to frequency error at smaller pulse widths.

The left-hand set of graphs can give us several more insights into the effect of pulse width on preshaping performance. The first thing that jumps out when we look at the perspective view of the surface is that the curve looks just like the one shown in Figure 3.8 (d) when the pulse width is near zero. This confirms our earlier claims that thin pulses can be treated as impulses. We have already discussed the increase in the residual at $\omega_N/\omega_d = 1$ as pulse width increases. What we have not yet noted is that the residual at $\omega_N/\omega_d = 2$ decreases at the same time. This is an artifact of the normalization scheme. As the pulse width increases, the total command size increases in magnitude, and the normalizing impulse increases with it. The maximum response to a single pulse does not grow, however. This means that when the pulse width is $T/3$ and $\omega_N/\omega_d = 9/8$, a situation which results in a single firing lasting for a period and a half, the impulse normalization tells us that the residual is only about 20% (see the dashed contour in the lower left graph of Figure 3.11). But we have fired for an odd number of half-periods, which we already know will produce the largest residual possible. Our only conclusion can be that the 20% figure is deceptive. Indeed, if we used the impulse normalization to look at the frequency robustness of the L^2 sequence we would find that as long as the ratio ω_N/ω_d is greater than one, the residual never climbs higher than 20%, a result which is clearly suspect.

It would seem that the only reasonable way of measuring the frequency robustness of the L^2 sequence at larger pulse widths is to use the maximum pulse response normalization employed in the right-hand graphs of Figure 3.11. When this normalization is used, we get the expected peak value of 1 at $\omega_N/\omega_d = 3/8$ and $\omega_N/\omega_d = 9/8$ for a pulse width of $T/3$. We also get to see a

strikingly different behavior at higher frequencies. At some frequency ratios and pulse widths the residual produced by an L^2 sequence is worse than could possibly be produced by a single pulse of any width. In these situations the preshaping sequence is actually amplifying the vibration. This is an excellent reason for keeping the pulse size as small as possible.

One last point needs to be made here. Even though the effect described above tends to drive us toward smaller pulse widths, we can see from the graphs that the greatest robustness to frequency error is not found at the smallest pulse widths. If we define an error limit beforehand we can take advantage of the 'V' shape of the valley around $\omega_N/\omega_d = 1$ and gain a little extra robustness. Looking back to the curve shown in Figure 3.10, it seems we can up the pulse width to about 12% of one period while still keeping the residual under 5% for either normalization scheme. Using a pulse width of this size gives us an increase of 20-30% in tolerable frequency error.

3.4.3 Time Cost

Measuring the time cost for an L^2 sequence is not much different than measuring it for a line sequence. The extra time required for the preshaping starts at one full period, the same as the time delay for the equivalent impulse sequence. As the sequence fills out, the delay shrinks, until the pulses start to run into each other and the minimum cost of $T/2$ is reached. When the command size increases beyond this point, however, we employ a slightly different strategy than was used with the line sequence. With the line sequence, we merely started up another sequence from the point where the first full sequence ended. If we did that with an L^2 sequence, the time cost

would grow by $T/2$ each time we restarted. To avoid this, we can intersperse the initial pulses of the new sequence with the final pulses from the old, as is shown in Figure 3.12. Except for small gaps near the half-period and period marks when the pulse width doesn't divide integrally into the period, we get a pattern of rapid on-off dithering at either end with a single, long pulse in-between. The time cost graph looks just like the one for the line sequence only shifted up by $T/2$.

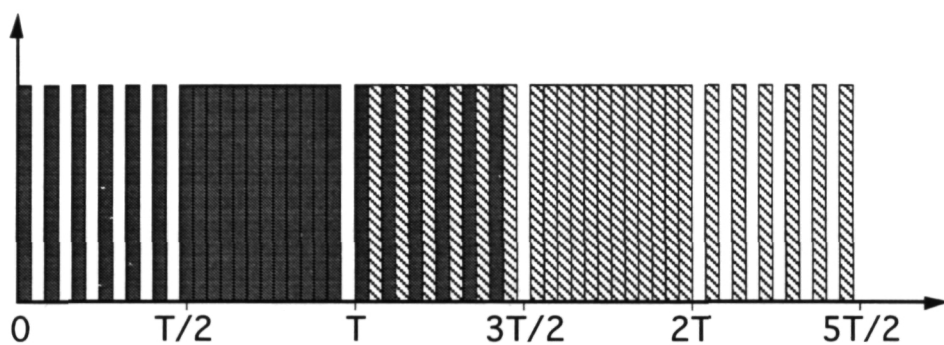


Figure 3.12: Concatenation of an L^2 sequence for large command sizes.

3.4.4 Resolution

As we start to look at more and more complex pulse sequences, the question of command resolution starts to come up. A line sequence takes $2n$ pulses to complete, while an L^2 sequence requires a total of $4n$ pulses. Real control systems will operate in discrete time, so there will be a minimum switching time and minimum pulse width which can be achieved by the control system. All of this means that when we use a preshaping sequence we sacrifice some level of fine control. If a line sequence is to be used, we cannot achieve commands which require an odd number of pulses. If an L^2 sequence is to be used, the resolution drops by another factor of two.

These limits, while they exist, are somewhat artificial constructs. In an ideal world we find that we cannot achieve zero residual vibration using a line sequence for command sizes which require an odd number of pulses. In the real world matters are substantially more fuzzy; we know that for most systems there is going to be some residual no matter what we do. Nonlinearities and frequency errors will take care of that. So if the command size is fairly large, we can execute it using a slightly shorter shaped sequence and then toss in the final pulse pretty much anywhere, figuring that the residual from that isolated minimum-size pulse is likely to be lost among other errors. For example, if we find that our command requires exactly 19 minimum-size pulses, we can achieve it in a variety of ways. We could ignore preshaping techniques and combine all 19 into a single large firing. We could use a line sequence and tack the odd pulse onto the first firing, giving us one firing of length 10 and another of length 9. Or we could use an L^2 sequence, tacking a line sequence onto the end of the first two sections and the last odd pulse after the first pulse in the line sequence. These three alternatives are illustrated in Figure 3.13.

Again, this mix and match technique is unlikely to cause much degradation in performance when the command size is large and we are comparing our residual to the maximum single-pulse response. When the command size is small and we are comparing to an equivalent single pulse firing, however, adding in the odd pulse or sequence fragment could cause a noticeable decrease in performance. The difficulties presented by the limited resolution of some sequences are partially addressed by the development of the new sequences that will be described in the next section.

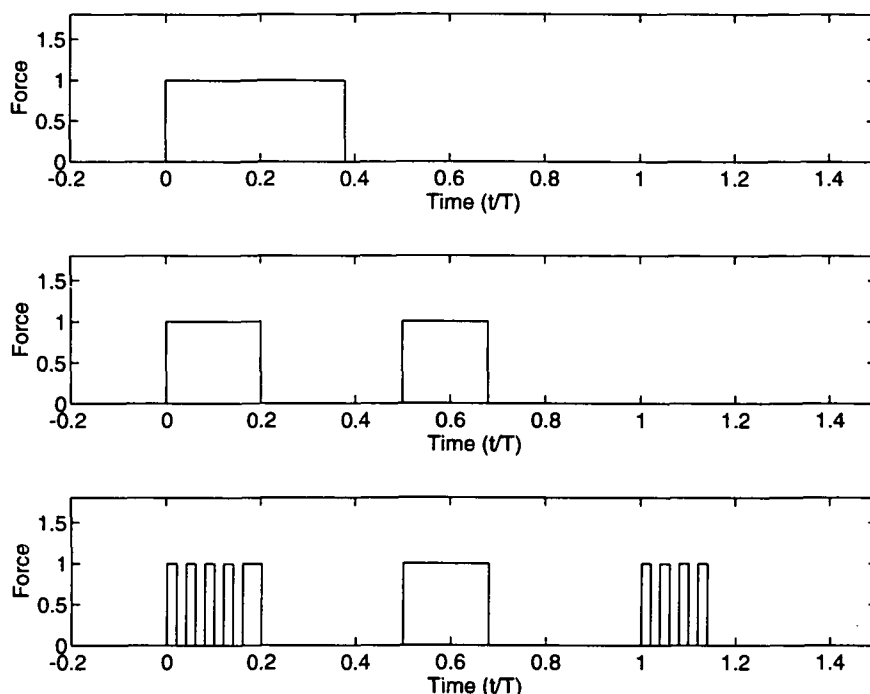


Figure 3.13: Force profiles for different methods of splitting up a total firing time of 19 times the minimum pulse width. At top is a single pulse; in the middle is a line sequence with the odd pulse tacked onto the first pulse of the sequence, and at bottom is an L^2 sequence with the three odd pulses attached as a line sequence and a single pulse.

3.5 Y and S Sequences

In all the work that has been done to date in command preshaping, a primary concern has been to find the sequence which meets the performance criteria with the minimum time delay. In Singer's thesis this was achieved by various convolutions of line sequences. [24] Rappole added a new dimension to this process with the symmetric sequence. [17] Singhose arrived at different sequences by relaxing the constraints somewhat at the design frequency. [25] Still, the library of available sequences remained fairly limited. In the constant amplitude domain, new performance criteria make the problem substantially more complex. Now insensitivity to frequency error must be

traded off against pulse size, number of firings, and command resolution as well as time cost. These new performance measures make attractive a new set of sequences which have not been examined in prior research.

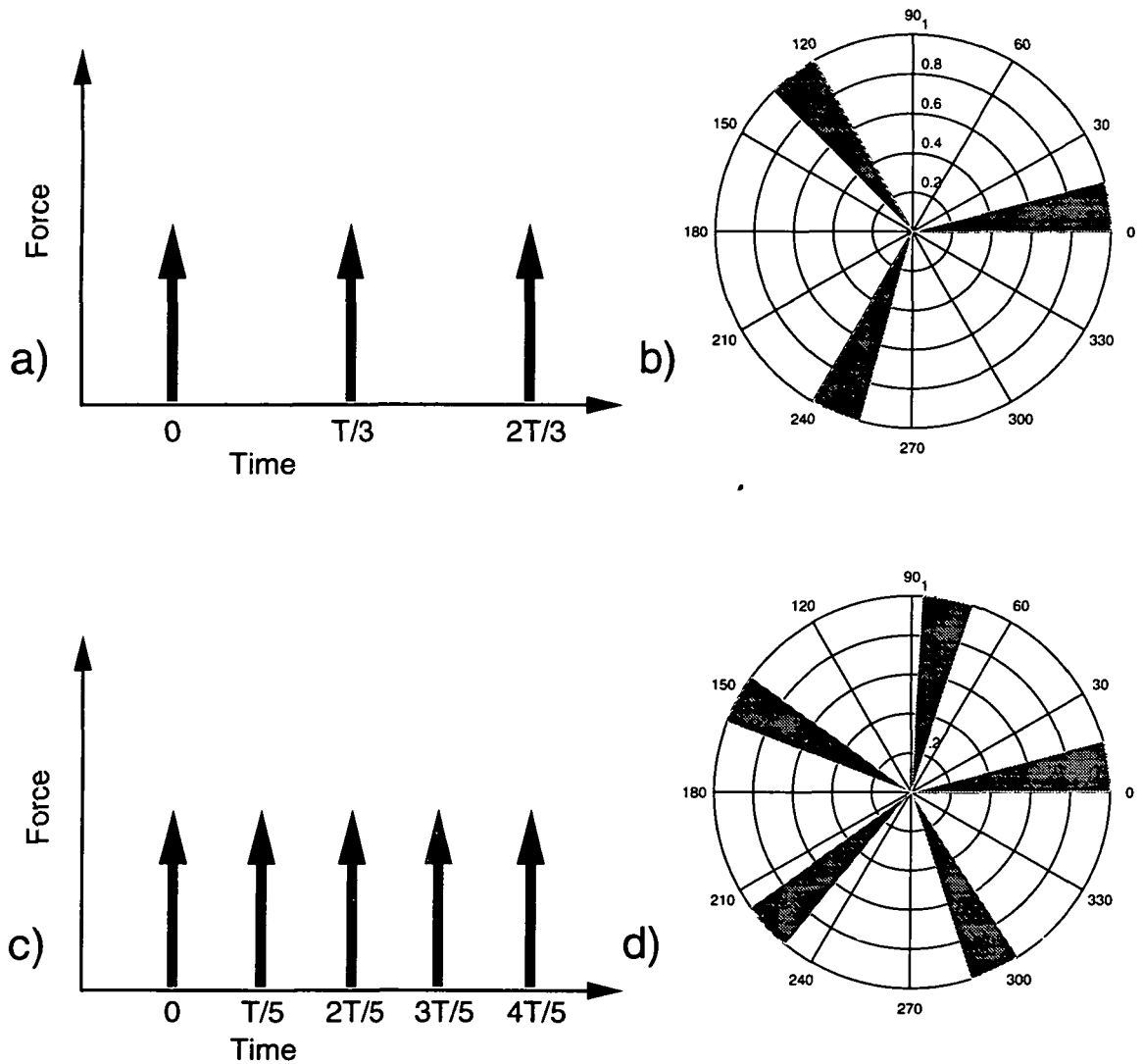


Figure 3.14: The Y and S sequences: a) impulse representation of a Y sequence; b) pulse sector diagram for a Y sequence; c) impulse representation of an S sequence; d) pulse sector diagram for an S sequence.

The two basic sequences to be discussed in this section are shown in Figure 3.14. The basic concept behind the sequences is clearly shown in the pulse

sector diagrams. Direct cancellation as in a line sequence is not the only way to achieve zero residual at the design frequency: dividing the command up into thirds to form a Y shape or into fifths to form a star (hence the S) shape will work just as well. Of course, these sequences do have an obvious disadvantage relative to the line sequence; they require respectively $2/3$ and $4/5$ of a period for completion. They also have advantages, however, which are perhaps not so obvious at first.

Figure 3.15 compares the insensitivity curves for L, Y, and S sequences. For normalized frequencies less than one, the Y and S sequences offer minimal improvement over the line sequence. It is at the higher frequencies that these new sequences show their worth. Whereas a line sequence produces zero residual for any odd frequency ratio, Y and S sequences touch down 33% and 60% more often, respectively. This shorter distance between frequencies with zero residual can be used to improve robustness to a single frequency or to shape for two different frequencies.

One particularly attractive sequence which uses this characteristic is obtained by convolving an L sequence with a Y sequence designed for a frequency one third lower than the L sequence. The resulting impulse sequence has pulses of magnitude 1, 2, 2, and 1, each separated by a half of a period for a total length of $1.5T$. The Y sequence keeps the residual magnitude down at its design frequency of $2/3$ and also at $4/3$, twice its design frequency, while the line sequence keeps the residual down at its design frequency of 1. The insensitivity curve of this combined sequence is shown in Figure 3.16 compared with an L^3 sequence (with pulses in a $1/3/3/1$ ratio), which requires

an equal amount of time for execution. If we admit a residual magnitude of 7.5% as acceptable, this LY combination sequence is robust to frequency errors of 40%, while the L^3 sequence is robust only to errors of 25%.

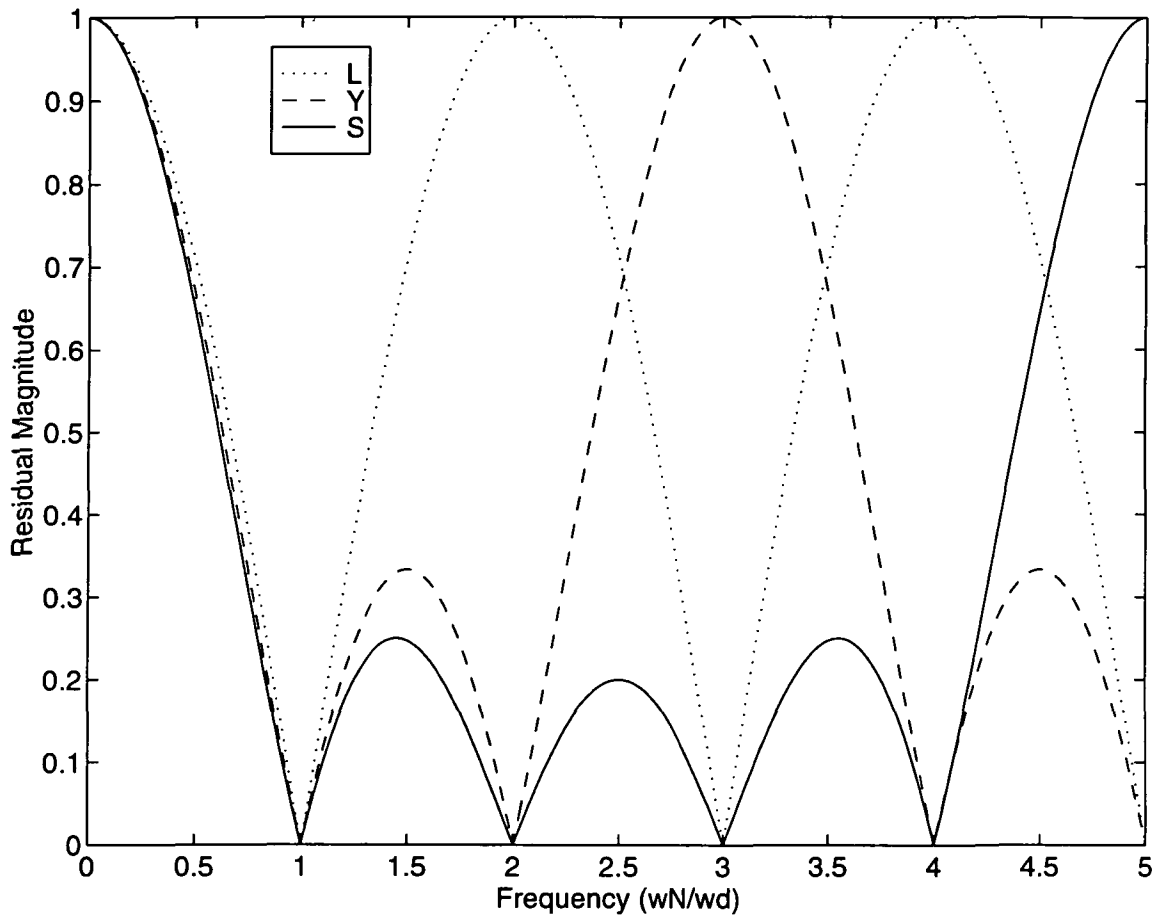


Figure 3.15: Insensitivity curves for L, Y, and S sequences. Note that the Y and S sequences touch down at the zero residual mark more often than the L sequence.

We can also take advantage of the Y sequence to gain robustness at a single frequency without having to approximate a double-amplitude impulse. We do this by convolving an L sequence with a Y sequence instead of another L sequence. By doing this we get the same (actually a little better) robustness near the design frequency and better performance at slightly higher frequencies as is shown in Figure 3.17. More importantly, however, we get an

LY sequence that has 6 separate pulses which are all the same magnitude! If we use a sequence of this sort no performance hit need be taken for the conversion from the continuous domain to the constant amplitude domain, and fewer firings will be needed to carry out the command. This is one of the most important benefits we get from using the longer base sequences.

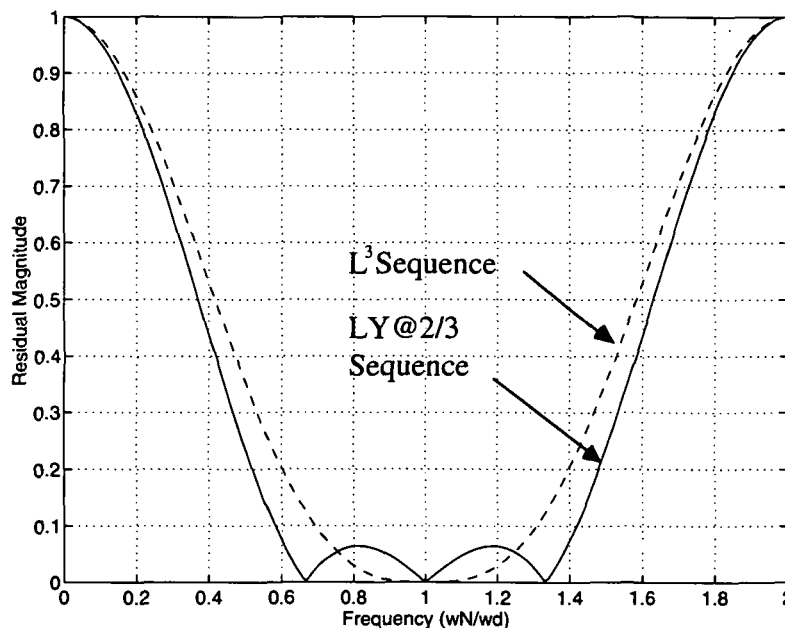


Figure 3.16: Insensitivity curves for the L^3 sequence and the LY@2/3 sequence. Note that though both sequences take the same amount of time for completion, the LY@2/3 sequence offers a substantially wider area of robustness.

In some situations, where system frequencies come in whole number ratios, use of these sequences may obviate the need for convolution of multiple line sequences and end up imposing a lower time cost. For example, if we have system frequencies at 1 and 2 Hz, we could either use a Y sequence or an L sequence convolved with another L sequence designed for the higher frequency. Both would give zero residual vibration at the system frequencies. The latter sequence would have a base time cost of $3T/4$, however, while the Y sequence would incur a time cost of at most $2T/3$.

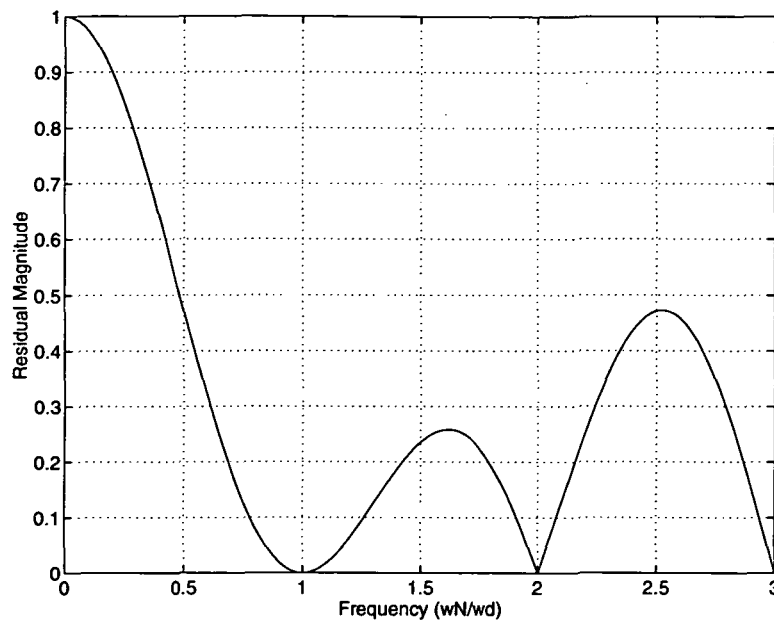


Figure 3.17: The insensitivity curve of an LY sequence. The robustness at the design frequency is equivalent to that of an L^2 sequence, but no impulses of differing magnitude were used.

The use of Y and S sequences can help us with command resolution as well. The odd number of divisions in both of these sequences complements the even number of divisions found in the L and L^2 sequences. For example, take the sample command used in the previous section, which required a total firing time of 19 times the minimum switching time. Before, we had to tack a line sequence and a stray pulse onto the end of the base L^2 sequence. By using a Y sequence instead, we can achieve the desired command size without having to introduce any uncanceled pulses. Similarly, a command of size 6 can be achieved by convolving an L sequence with a Y sequence instead of adding an L sequence onto the end of an L^2 sequence. Whereas the latter alternative would sacrifice some performance, the LY sequence would actually be more robust to frequency error than a normal L^2 sequence.

The Y sequence has one more important characteristic. Its $T/3$ spacing is just what is needed to construct pulse sequences that minimize the loads (the force in the spring in Figure 3.1) induced in the system. The next section uses the Y sequence as a base for constructing a new set of sequences which minimize the loads induced by system flexibility.

3.6 Load Sequences

So far we have concerned ourselves only with preventing unwanted motion in the system of interest. We have measured our success in terms of the amplitude of the vibration remaining after the firing sequence has been completed. In many cases, however, the main concern is not with the flexible motion of the system but the loads experienced at various locations. Often design specifications will require that the loads on a part not exceed a given limit. In this section we will describe a type of pulse sequence that will keep the loads in a flexible system below a given level while imparting the required impulse in a minimum amount of time.

The first pulse in a command sequence will cause the load to oscillate in a sinusoidal fashion. The amplitude of this sinusoid sets a lower bound on the maximum load to be experienced by the system. Depending on the timing of the other firings in the sequence, the maximum load may exceed this level, but as long as only positive pulses are allowed the maximum load will be at least this high. The amplitude of the sinusoid set up by this initial firing is of course determined by the impulse the firing imparts, which is in turn determined by the length of the pulse. So to create our minimum-load

sequence we select a pulse width appropriate for the desired maximum load and then make sure that the remaining pulses do not cause the response to exceed that limit.

By expressing the amplitude of the residual vibration as a function of the spacing between pulses, we can find the spacing necessary to keep that amplitude from exceeding the limit set by the initial pulse. Approximating the pulses as impulses, we can say that the residual from two pulses, one starting at $t = 0$ and the other at $t = t_1$, is

$$\chi(t) = \cos \omega t + \cos \omega(t - t_1) \quad (3.7)$$

which can also be expressed as a sinusoid with amplitude

$$A = \sqrt{2(1 + \cos \omega t_1)} \quad (3.8)$$

and phase

$$\phi = \tan^{-1} \left(\frac{-\sin \omega t_1}{1 + \cos \omega t_1} \right). \quad (3.9)$$

The minimum value of t_1 for which the amplitude is equal to one is $t_1 = T/3$. Substituting this into equation 3.9 we find that the phase is equal to $-\pi/3$ ($= -T/6$). The next pulse would then come at $(T/3 - T/6) + T/3 = T/2$. To continue the pattern we can keep adding pulses spaced by $T/6$; to finish it off leaving zero residual vibration we can add a final pulse set off by $T/3$ seconds. All in all, we fire a single pulse at t_0 followed by a sequence of n pulses beginning at $t = t_0 + T/3$ and separated by $\Delta T = T/6$, finishing with a single pulse at $t = t_0 + 2T/3 + nT/6$. When $n = 0$ we have a simple Y sequence. When $n = 1$ we have the sequence shown in Figure 3.18. In the first graph the dotted line shows the force profile of a single pulse of equal total impulse; in the

middle two graphs the dotted line shows the behavior induced by such a pulse in the time domain. From the response curve we can see that the load sequence does indeed produce zero residual vibration, while from the loads curve we can see that the force felt by the second mass never exceeds the peak from the first pulse. In the limit as n goes to infinity the loads curve will flatten out and approach a square wave.

These load sequences provide a varying degree of robustness to uncertainties in system frequencies. A load sequence constructed using the algorithm described above will be robust to underestimation of the system frequency, but not overestimation. This can be seen in the bottom right-hand graph of Figure 3.18, which shows how the maximum load varies with the ratio of the system's actual frequency to the design frequency. In this particular case, the maximum load will be minimized for ratios between 1 and 1.5. Robustness in both directions may be achieved by shifting the sequence such that its high performance region is centered about a ratio of one. For the case where $n = 1$ shown in Figure 3.18, this would be achieved by multiplying each pulse time by $5/4$, giving a modified sequence with pulses at $t = 0, \frac{10T}{24}, \frac{15T}{24}, \frac{25T}{24}$. The range of frequencies over which the loads can be minimized is quite wide for $n = 0$ but shrinks by a factor of two for each additional pulse added. The sequence shown in the figure was chosen because it is robust with respect to both residual vibration and loads in the same frequency range.

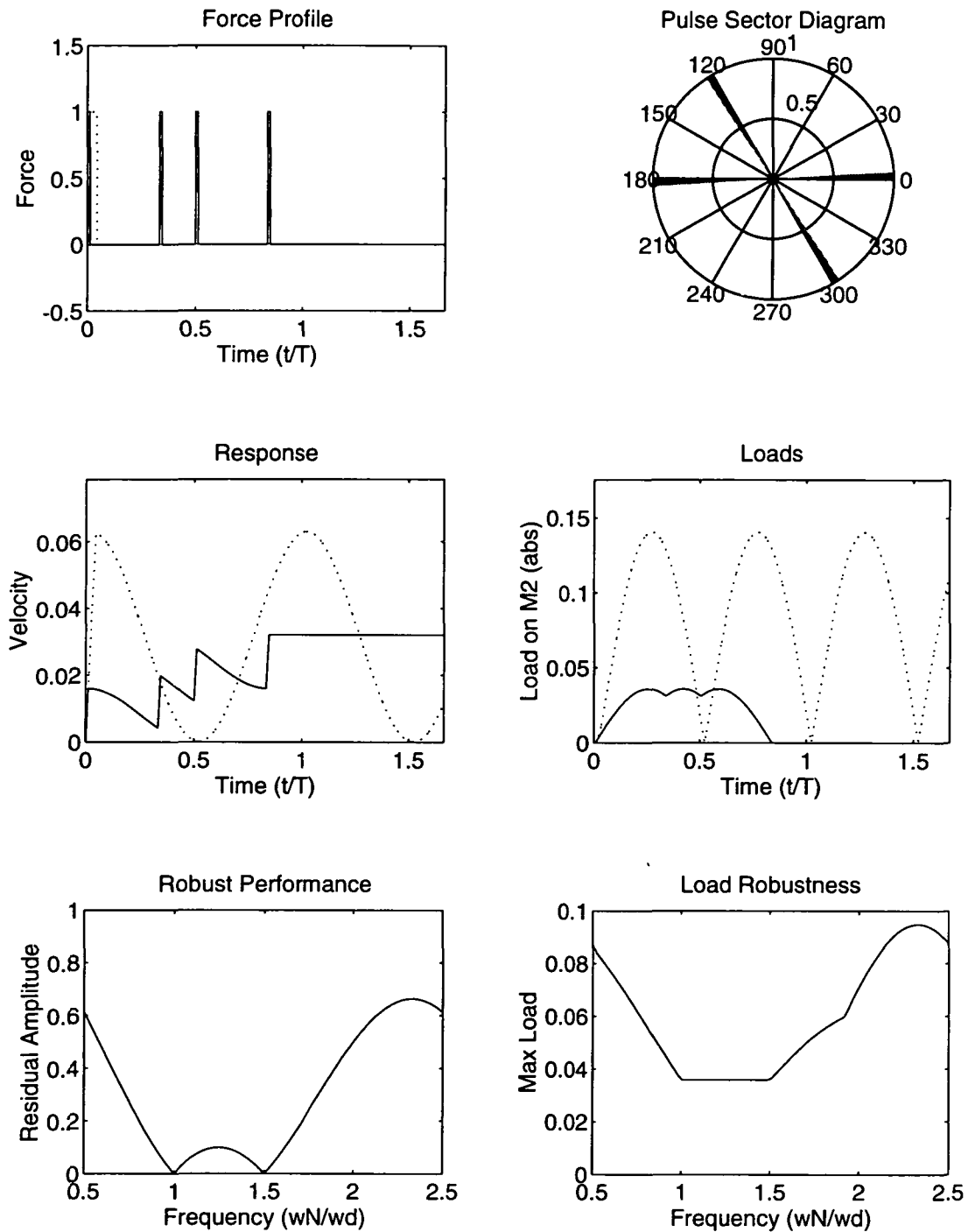


Figure 3.18: Load sequence for $n = 1$. At upper left we have force profiles for the sequence and an equivalent single pulse. At upper right, we have a pulse sector diagram for the sequence. The middle two graphs show time histories for the load sequence (solid) and the single pulse (dotted). The graph at bottom left shows the insensitivity curve of the sequence, and the one at bottom right shows the variation of the maximum load experienced during the move with frequency.

The robustness of these load sequences with regard to residual vibration is dependent on the length of the sequence in a much more complex way. For any value of n there will be no residual vibration at the design frequency, but the robustness in that area and the residual at other frequency ratios change dramatically as n increases. Figure 3.19 shows the evolution of the insensitivity curve as new pulses are added to a load sequence. In the graph the residual magnitude is indicated by color, with small residuals appearing as red and residuals near unity appearing as magenta [In the black and white version, both the red and magenta show up as darker shades of gray. Since the magenta regions are restricted to the top, bottom, and left sides of the graph, the information content of the graph should remain intact.] We immediately notice an overall trend toward the red (low residual) end of the spectrum with longer sequences. The next feature that pops out are the widening bands of red at frequency ratios of 1 and 5. These show that the robustness around the design frequency increases with the number of pulses. Of course, the most striking feature of the graph is the curves which the red patches seem to make, like a knit fabric or two intersecting wave fronts. Unfortunately, though it seems there must be some meaning to such a clear pattern, no explanation for the behavior has yet been forthcoming. One may note, however, that for all sequences having an even number of pulses there is a region of low residual at a frequency ratio of 3, and for all sequences with a number of pulses divisible by three there are regions of low residual at frequency ratios of 2 and 4.

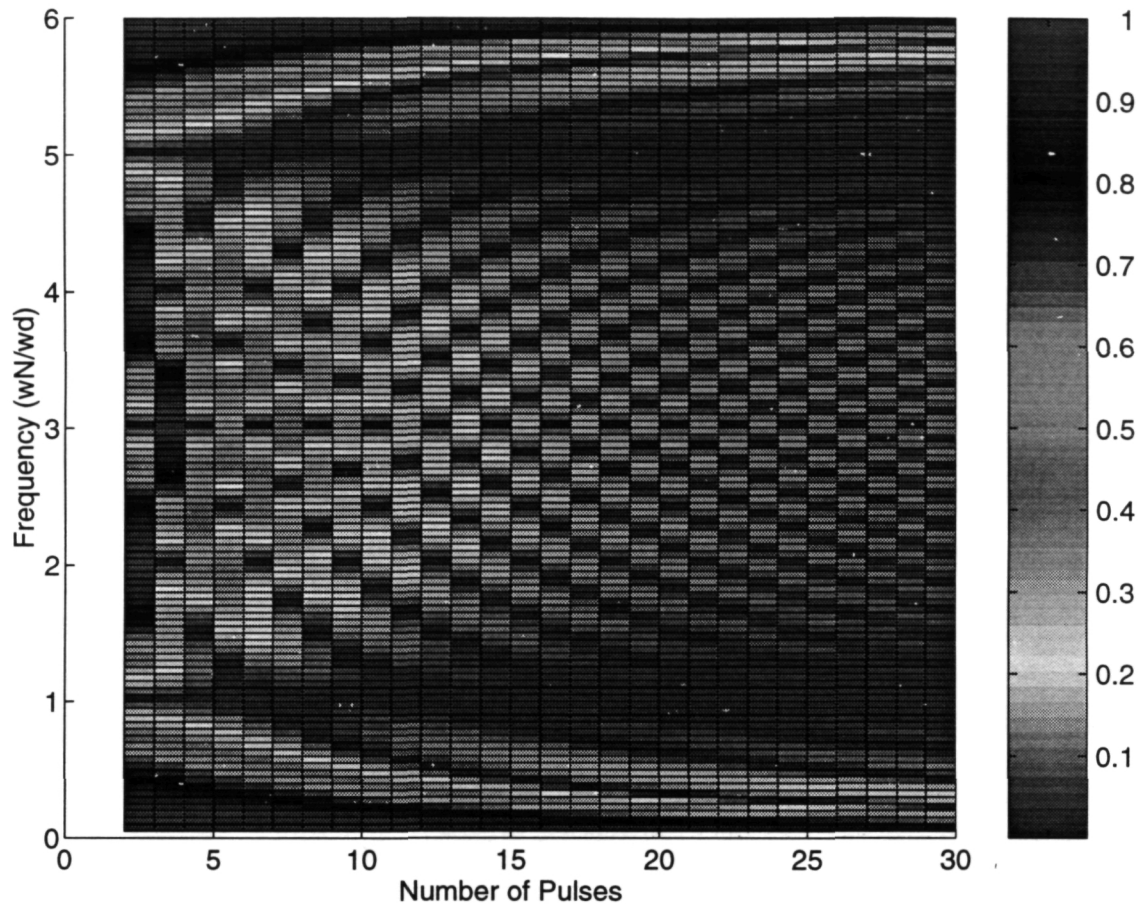


Figure 3.19: Map of the effect of sequence length on insensitivity for load sequences. Each vertical strip represents the insensitivity curve for a sequence of the specified number of pulses viewed from above. Low residual levels are indicated by red (dark gray), while high residuals are marked in magenta (which also shows up as dark gray in the B&W version). The regions of high residual are present only at frequencies near zero or six and at low numbers of pulses in the interior region. These latter areas can be identified by looking at the insensitivity curves for the appropriate sequences shown elsewhere in this thesis.

By looking only at how the residual vibration's robustness to frequency error changes with increasing sequence length we ignore an important point, however. As has been said before, the loads robustness to frequency error is halved for each additional pulse which is added. By the time a sequence gets to the right-hand side of Figure 3.19 very little loads robustness will remain.

An alternative way of generating large commands with minimal loads and high robustness to frequency error for both loads and residual vibration would be to use an $n = 1$ load sequence as a building block for constructing larger sequences. As soon as one sequence is completed, another can be begun. Sequences constructed in this manner will take a sixth of a period longer to complete for each group of four pulses, but the gain in loads robustness may be worth it.

3.7 More Complex Systems

Although the model proposed in Section 3.1 is excellent for capturing the fundamental issues of constant amplitude preshaping, it neglects or ignores several important properties of real systems. Unlike the model, real systems have damping. Real digital control systems also operate in discrete time increments. Finally, real systems often will have more than one significant mode of vibration. All of these properties will affect the ultimate performance of a preshaping sequence; this section will strive to show the nature of these effects, their extent, and ways of compensating for them.

We will begin with a look at damping. In most applications where command preshaping is considered as an option, damping will be low for the simple reason that if it is not, preshaping becomes unnecessary. Due to the absence of external viscous forces, many space systems have low damping ratios. The space shuttle RMS, which serves as a primary example in this thesis, has a structural damping of 2.2%. When friction and viscous damping

are added in, it may get up to 4% or so. This is relatively small, but it is not zero, so neglecting it will have some effect.

The impulse sequences presented by Singer do take into account damping. [24] The later impulses in the sequences are simply reduced in magnitude so that they will properly cancel the decayed response from previous impulses and shifted in time to go by the damped natural frequency instead of the natural frequency. In a line sequence the second pulse gets multiplied by a factor of K and delayed by a factor of γ , where

$$K = e^{-\frac{\zeta\pi}{\sqrt{1-\zeta^2}}}$$

and

$$\gamma = \frac{1}{\sqrt{1-\zeta^2}}.$$

In an L^2 sequence the impulses come in a ratio of 1 to $2K$ to K^2 . It is when we look at the value of K that the difficulty of coping with damping becomes apparent. For a 5% damping ratio, K is about 0.854. This means that for an L^2 sequence the three impulses will come in a ratio of 1 to 1.709 to 0.73, which is not even remotely close to a whole number ratio.

The importance of pulse amplitudes which are related in whole number ratios cannot be overrated. This is because real systems have a definite lower limit on pulse sizes, and all larger pulses must be constructed from integer combinations of these smallest building blocks. So to produce a reasonable facsimile of an L^2 sequence designed for a 5% damping ratio, we would have to use a pulse of width 10 followed by a pulse of width 17 and another of width 7. Even this is about 5% off for the last pulse. In addition to its low resolution —

all commands must come in multiples of 34 times the minimum pulse width — this sequence would likely start running into the difficulties which come when pulse widths become significant compared to the system period. Differences in spacing would also make the interleaving described in Section 3.4.1 (Figure 3.9) impossible. Overall, this seems like a poor solution.

Of course, not all damping ratios will be as intractable as the one given. Figure 3.20 shows how the relative pulse sizes change with damping ratio for an L^2 sequence. For a damping ratio of about 22%, the first and second pulses are equal in magnitude and the last pulse is one fourth the size of the first two. Admittedly, if the damping ratio is that high the use of impulse preshaping becomes rather questionable. Yet similar points may exist in other sequences at more realistic damping ratios, so the possibility of such a lucky coincidence should be considered. In fact, for the RMS's damping ratio of 4%, the ratios work out rather well, so pulses in a ratio of 4 to 7 to 3 could be used with relatively little error.

By this time we can see that taking a system's damping into account requires quite a bit of extra work. It is only natural to wonder if the performance increase gained by accounting for damping is worth the trouble. The answer to this question is "probably not." Figure 3.21 below shows the effect of unmodeled damping on the magnitude of the residual vibration for three different sequences. For a simple line sequence the error can become significant if damping ratios higher than a few percent are neglected. Line sequences are rarely used by themselves, however, and if an L^2 sequence is used the error from neglecting a damping ratio as high as 5% is only about a

half of a percent. It appears that for most sequences our assumption that damping may be safely ignored is quite justified.

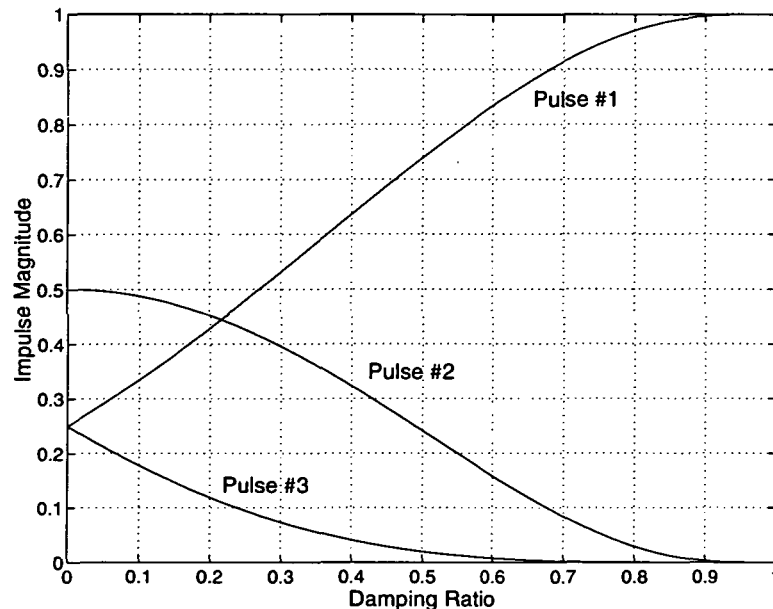


Figure 3.20 Effect of damping ratio on impulse magnitude for an L^2 sequence.

Damping is not the only factor which keeps us from achieving optimal performance with impulse preshaping. Another factor which will add in a small but noticeable error is the discretization of time in digital controllers. In most control systems pulses can only be begun at discrete time intervals, which will only rarely coincide with the desired firing times to high precision. This slight mislocation of the pulses in the sequence will make it nearly impossible to get a near-perfect cancellation. The larger the time step which the controller operates under, the larger the residual will be.

Finally, we know that in real systems there will often be more than one significant mode of vibration. If we are not careful, we may find that a pulse sequence which reduces vibration at one frequency amplifies it at another.

The basic design tool for coping with vibration at multiple frequencies is the insensitivity curve; we seek a sequence which robustly approaches zero residual vibration at each of the design frequencies. As we shall see, there are a variety of means for achieving this end.

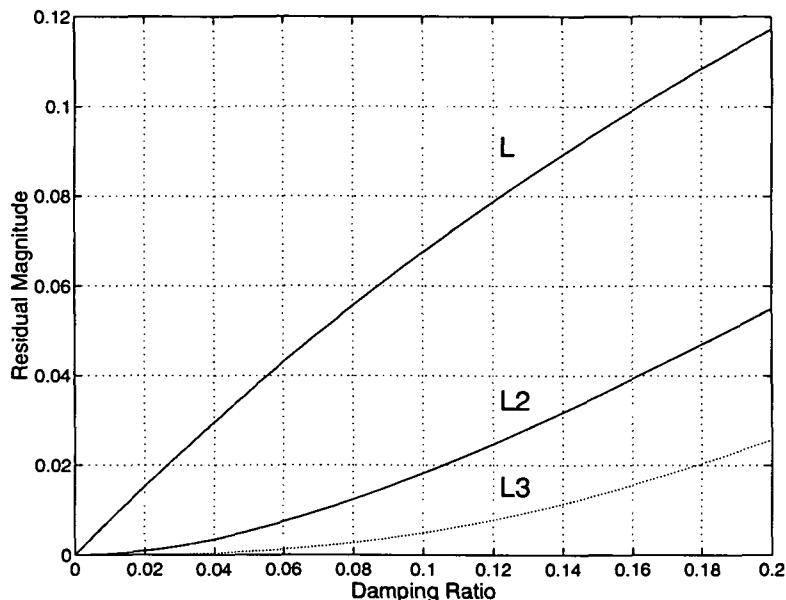


Figure 3.21: Error induced by the assumption of zero damping as a function of actual damping ratio for three different sequences. The numbers shown assume that the design frequency and natural frequency are the same.

When we convolve two sequences, we multiply their insensitivity curves on a point by point basis. This is why the L^2 sequence is more robust than the L sequence; by convolving an L sequence with itself we square the magnitude of the residual at each frequency. If we convolve an L sequence with itself twice we get an L^3 sequence, which is even more robust. We can also use this property with different types of sequences. Figure 3.22 shows the insensitivity curves of an L and a Y sequence and then the insensitivity curve of their convolution, an LY sequence (which also happens to be a minimum load sequence).

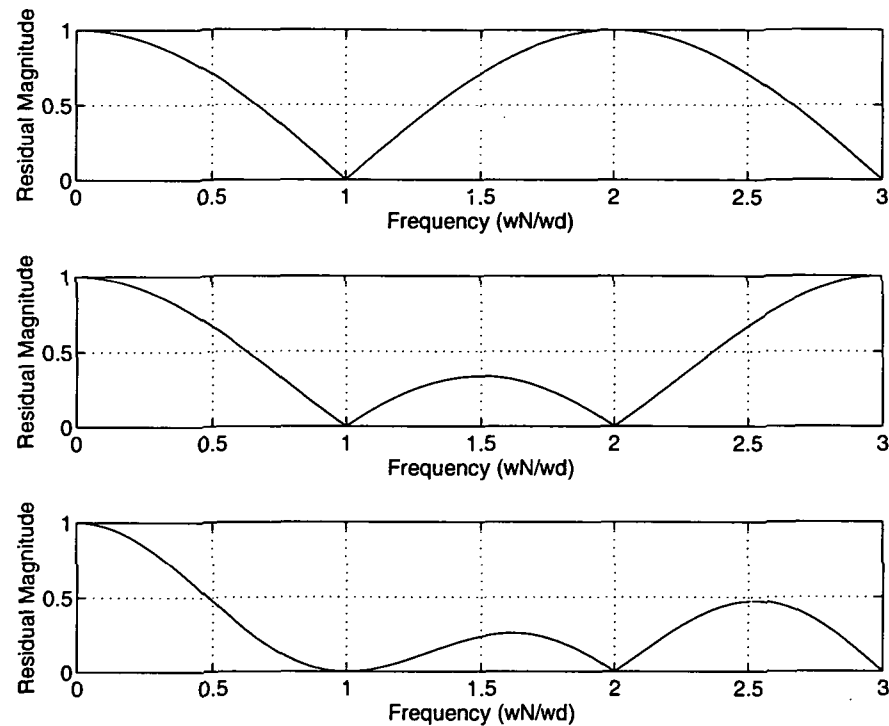


Figure 3.22: Illustration of the effects of convolution on sequence insensitivity. The top curve is the insensitivity of an L sequence, the middle curve shows the insensitivity curve for a Y sequence, and the bottom shows the insensitivity of their convolution, an LY sequence.

Thus the simplest way of creating a sequence to prevent vibration at two frequencies is to create sequences for each frequency separately and then convolve them. If extra robustness is needed at one or both frequencies just convolve in extra sequences as appropriate. Unfortunately, there is one problem with this method: the time cost of the preshaper will increase rapidly as each additional sequence is convolved in.

If we are lucky, we may find that some of the modes we are trying to cancel come at frequencies which are near whole integer multiples of each other. In these cases we can utilize base sequences with valleys in their insensitivity curves that fall naturally near the frequencies in question. For instance, if

one system mode comes at roughly twice the frequency of another we can incur a smaller time cost by using a Y sequence instead of convolving two line sequences.

A third way of coping with multiple modes would be to use the symmetric sequences developed by Rappole. [17] These sequences incur a smaller time cost than sequences formed by convolution, but in general have poorer insensitivity curves. Another serious disadvantage of symmetric sequences in a constant amplitude framework is that the pulses in the sequence come in highly non-integer ratios. Here we get the same problems that we encountered earlier in dealing with damping: loss of command resolution and increased time cost for larger commands.

3.8 Designing a Constant Amplitude Preshaper

In this chapter we have discussed a wide variety of issues relating to impulse preshaping in constant amplitude systems. A number of different sequences have been presented, each of which has its own unique advantages and disadvantages. The question which remains to be answered is how to select one of these sequences and conduct the necessary trade-offs for a particular application.

The first things that need to be considered are the physical constraints of the actuators which drive the control system. We need to know the minimum pulse width that the actuators can generate as well as the effect of rapid switching on actuator system wear. If the actuator design specifications limit the number of switches to be made in a given time period, that should be noted

here also. The velocity increment from a minimum-width pulse should be compared to the desired command precision to give an idea of what resolution is needed. If the actuators are extremely powerful we may have only a few minimum-width pulses to play with and hence little flexibility in sequence design. However, if the actuators have relatively little control authority a longer burn time will be necessary and we will have greater flexibility in sequence design.

Next we need to think about the frequencies at which we intend to prevent vibration. Determining what those frequencies are is a whole separate issue which will be discussed in the next chapter. Here we assume that the frequencies which produce the largest vibration have already been identified to some degree of accuracy. Once the frequencies have been given, a quick check should be made to see if any of them occur at near integer multiples of each other. If one or more does, use of a complex sequence such as a Y or S sequence may be a time saver. Otherwise, we should probably base our sequence on some convolution of line sequences.

Along with each frequency to be removed we need an assessment of the frequency uncertainty and a specification of the maximum allowable residual. The former quantity will give us an idea of how many convolutions we need at each frequency. The latter quantity shows us how much room we have to modify the base sequence. One thing this will affect is the decision on what pulse width is to be used. We can gain some benefit in frequency insensitivity by using larger pulse widths as long as those widths remain relatively small compared to the smallest period being shaped for. This is illustrated in Figure

3.11. Of course, the benefit that can be extracted from this may be limited by the growth of the residual at the design frequency, which is illustrated in Figure 3.10. This growth occurs rather slowly, though, so in many cases the limiting factor will be the need to keep the pulse width small compared to the period, not the maximum allowed residual. In these cases we may obtain further robustness by separating the design frequencies of the line sequences to be convolved. For example, we could convolve one line sequence at a frequency of 0.9 times the design frequency with another at 1.1. As this gap increases the robustness to frequency error improves, but the residual at the design frequency also increases. When the residual at the design frequency hits the maximum allowed value we get the greatest insensitivity to frequency error.

Sadly, there is not always a hard and fast "maximum allowed residual" available. Often, we have to trade off a sequence's robustness to frequency error with the time cost it incurs. Using an L^2 sequence instead of an L sequence triples the average time cost (The average time cost is $T/4$ for an L sequence and $3T/4$ for an L^2 sequence) but provides greatly improved insensitivity to frequency error. Here is the place where the most "engineering judgment" comes into play. Using whatever method, an acceptable time cost must be determined and a sequence selected.

Having selected a sequence and a pulse width, we then proceed to construct individual commands. We do this by concatenating an appropriate number of pulse sequences together. If the command cannot be achieved by concatenating an integer number these sequences, we make up the odd

fractions with either sequences of smaller width, simpler sequences (L instead of L^2 , for instance), or a combination of the two. The result may not look much like what would be expected from an impulse preshaper, but it should produce excellent performance.

The process described above is fairly involved, but can easily be streamlined into an automatic process. Depending on the application, many of the above steps could be taken out to make the process even simpler. Using it as a guideline for constructing sequences should be much faster than using a nonlinear optimizer to generate the sequences and provide a much greater understanding of exactly what is going on. Of course, there is always a place for optimization techniques. If the system in question has clear-cut, unchanging dynamics, a nonlinear optimizer can probably offer a slightly superior solution at a one-time computational cost.

Application to the Shuttle

Chapter 4

Now that we have gone over the theory of constant amplitude command preshaping, it is time to discuss the practice. This chapter will present simulation results showing how CAP techniques can be used with the shuttle RCS system to reduce vibration and loading in the RMS and its payloads. To demonstrate the applicability of the techniques to the entire range of shuttle-RMS operations, the simulations will show the techniques in operation with a variety of different payloads, including the Upper Atmospheric Research Satellite (UARS), the Hubble Space Telescope (HST), and a (now somewhat outdated) space station model. For each of these payloads CAP techniques will be applied to both vernier and primary RCS operations, with the performance of the preshaped sequences compared to that of single pulse firings and, for the primaries, with the alt-mode. A quick look will also be taken at application of command preshaping to an unloaded arm.

Before we present these results, however, a little background material is in order to establish the legitimacy of the simulation results. For that purpose, the first two sections describe the simulation used to model the RMS behavior and the particular way in which it was configured for the tests which were conducted. After that comes a section on frequency identification to show how the sequences used in the results section were selected. With this material to assure their accuracy, the results which follow should speak for themselves.

4.1 The DRS

There are quite a few simulations of the RMS floating around the space community these days. Johnson Space Flight Center has the Shuttle Engineering Simulation (SES), which is linked to real hardware and runs in real-time so that it can be used to develop real-time procedures. Testing of software changes is done using actual flight hardware with the SAIL simulation. SPAR, the company responsible for the original development of the RMS, has its own simulations, among which the one known as ASAD (All Singing All Dancing) is most prominent. Draper Lab makes use of several different RMS simulations, including LSAD (Less Singing and Dancing), the FRS (Fast RMS Simulator), and the DRS (Draper Remote Manipulator Simulation). Developed over a period of ten years, designed to account for the whole gamut of nonlinear effects, and verified using actual flight data, the DRS represents an extremely good model of the actual flight system.

Over the years, the primary use of the DRS has been for analysis of shuttle payload deployment and retrieval operations. Draper has been responsible for

predicting how the RMS will respond to planned flight conditions and selecting appropriate autopilot configurations to avoid instabilities or excessive payload loading. The DRS was validated for this type of work by comparison to actual flight data in the period from 1984 to 1985. In 1987 a flexible payload module was added to allow higher fidelity simulation of the internal motion of payloads with fragile flexible components such as solar arrays.

The code for the DRS was written in FORTRAN for an IBM mainframe, but has also been ported to a variety of different UNIX machines including Sun Microsystems and Silicon Graphics workstations. Though it does not run in real-time in its highest fidelity mode, it can approach real-time on some machines if some of the nonlinear options are turned off. The mainframe version of the DRS supports a connection to an autopilot emulator so that interactions between the orbiter's control system and RMS operations can be studied. On the UNIX side, the DRS has been linked to the Interactive On-orbit Simulation (IOS) to provide the same functionality.

One of the more important features of the DRS for our purposes is the capability to model flexible payloads. The standard DRS can model a flexible payload with up to 10 modes of vibration; plenty for most payloads which have flown so far. Recently, however, with the work done to support the HST servicing mission and future space station operations, complex payload models have necessitated an upgrade of the DRS flexible payload capability to handle more modes. Versions now exist that can handle as many as 50 flexible payload modes.

To get an idea of the level of accuracy which can be obtained by DRS simulations, one need merely look at the plots shown in Figure 4.1. The plots show the position of the end of the arm in POR (Point Of Resolution, a coordinate axis system which usually coincides with the payload axis system) coordinates while the RMS is moved from its hover configuration to the capture/release position during STS-61, the Hubble Space Telescope servicing mission. Though it might not be apparent at first, each of the three graphs contains two curves: the one marked with 'o' symbols is the output from a DRS simulation, while the one marked with '+' symbols is taken from actual flight data. With a simulation of this high quality we can have quite a bit of confidence in the accuracy of our results.

4.2 DRS Configuration

The DRS has a large number of options which affect its performance and the simulation quality. It is important to state up front which options were used and which were not in the simulations that were run. In general, the options were the same as those used in the actual analysis conducted by Draper in support of each mission.

In each of the simulations run for this paper, the main arm dynamics and joint servos were integrated at a rate of 1000 Hz. Nonlinear torques in the arm dynamics were updated at 25 Hz, and jet firings and simulation output were processed at the standard 12.5 Hz (80 msec timestep) rate.

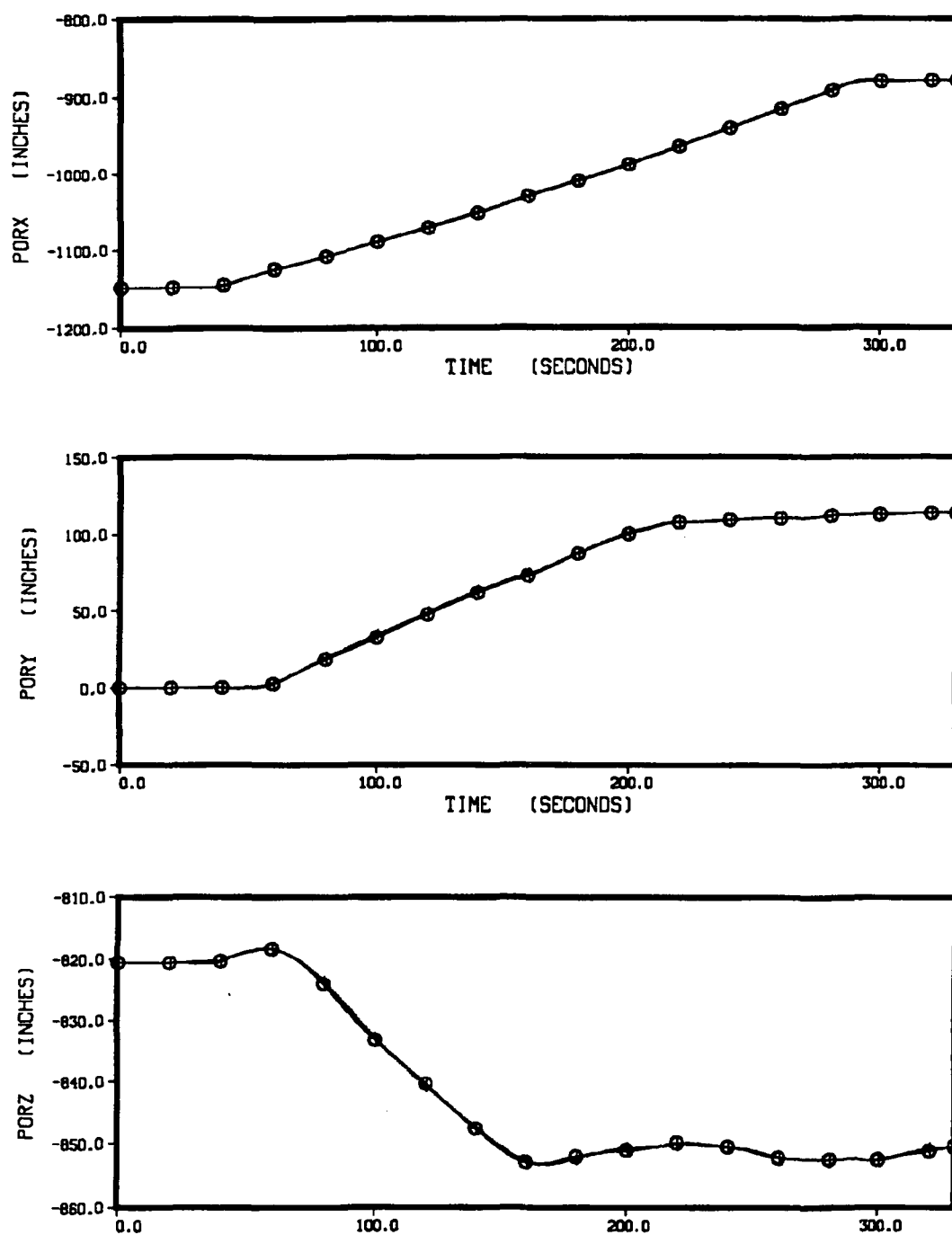


Figure 4.1: Comparison of results from a DRS simulation with actual flight data from STS-61, the Hubble Space Telescope servicing mission. The curve marked with '+' symbols represents the flight data, while the curve marked with 'o' symbols represents DRS output.

All of the simulations were run with RMS brakes on instead of using position hold, as a test of the theory with no outside interference was desired. The simulations using the Hubble Space Telescope (HST) as the payload made use of a more rigorous but more computationally intensive joint friction model than for the other payloads, which used a "hairbrush" friction model. The HST simulations also made use of a version of the DRS which could handle 23 flexible modes; the others used only 10. Lumping of arm flexibility at the grapple fixture was employed in all the simulations.

4.3 Frequency Identification

In order to apply the constant amplitude command preshaping techniques developed in the previous chapter, we need some way of determining the principal frequencies at which the system vibrates. In a simple mechanical system this can often be achieved by straightforward structural analysis. In a more complex nonlinear system such as the space shuttle with RMS and payload, it becomes a much more difficult task. Frequency identification in such a system can become a painstaking process yielding results that can easily vary by 10-20% based on the particular conditions under which the evaluation is conducted.

The frequency characteristics of the shuttle arm are dependent both on the mass of the payload being carried and the position of the arm within its workspace. The location of the fundamental frequency can vary by as much as an order of magnitude when comparing small to mid-sized payloads with something as heavy as the planned U.S. space station. Variations in arm

position can also produce significant changes in the arm frequencies, though not on the order of those produced by large mass differentials. A characterization of the dependence of frequency on arm position for an unloaded arm can be found in Singer. [24]

The internal vibrations of the arm are not all we need concern ourselves with, however. We must also watch the vibrational modes of the payload. These may occur at frequencies near those internal to the arm, thus raising the possibility of unpleasant resonant effects, or at entirely separate frequencies. It is usually these modes which will determine if the loads on a satellite's solar arrays are too great, so they are of critical importance.

A variety of different techniques exist for frequency identification, each with its own advantages and disadvantages. Three of the most commonly used methods are structural analysis, singular value decomposition, and spectral decomposition using fourier transforms. Each of these has its place in the analysis of flexible interactions in the shuttle-RMS system.

Structural analysis is used for the most involved frequency estimates and for the roughest. Finite element models provide the basis for simulation of payload dynamics and give detailed information about the natural modes of vibration internal to the payload. At the same time, engineers often use simple statistics such as the payload mass and position of the arm's end effector to get a rough estimate of the value of the system's fundamental mode. The payload dynamics predicted by the finite element model change when the payload is attached to the RMS and shuttle, however, and the estimates

generated from mass and position data are not precise enough for use in serious analysis.

Singular value analysis has the advantage of being able to determine modal directions as well as amplitudes. This helps with formulation of problems in terms of worst case scenarios. However, it assumes sinusoidal inputs (a far cry from the square pulses put out by the shuttle RCS system), and it bases its predictions on a linear model of what in our case is a complex nonlinear system.

The most frequently used tool for the task of frequency identification is probably the fast fourier transform, or FFT. The FFT allows us to take a discrete time history of a system's response to an input and extract the frequency content in what is known as a power spectral density (PSD). The PSD is a powerful tool, and the one which has been selected for use in this thesis. It must be employed with a certain degree of caution, however, as it only provides information about the frequency content of a signal; there is no directional information such as that provided by a singular value decomposition. For example, Figure 4.2 (a) shows a PSD of the response at the end of the arm to the firing of three jets in the -Pitch, -Roll direction. Figure 4.2 (b) shows a PSD of a similar time history, in this case in response to a maneuver in the +Yaw direction. Note that although the frequency contents are similar, they are not exactly the same.

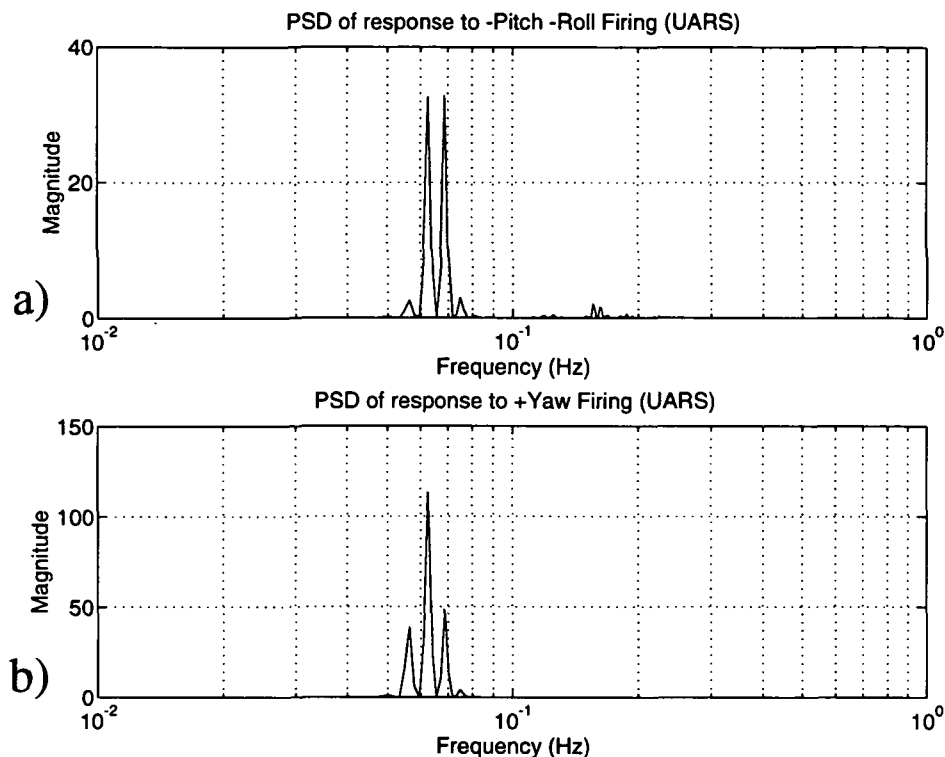


Figure 4.2: Power spectral densities of RSS POR position from simulation of PRCS jet firings in two different directions. a) PSD for -Pitch, -Roll firing; b) PSD for +Yaw firing.

The PSD shown in Figure 4.2 (a) was generated from the time history shown in Figure 4.3. The data plotted in the figure is the root sum square of the POR X, Y, and Z coordinates (POR stands for Point Of Resolution, an arbitrarily defined point in space relative to the tip of the arm, usually selected to bear some relation to the payload position). To generate the graph shown in Figure 4.2 (a), several operations must be applied to the data set. First of all, one may note that the graph in Figure 4.3 begins with a value of a little under 1088 inches, but that the mean value differs from this initial condition by about two inches. This differential is caused by brake slippage in the arm's joints occurring during the first few seconds when loads are greatest. These first few seconds of the data set are discarded to avoid

distorting the PSD with this effect. Next, the signal's mean is calculated and subtracted off in order to remove the DC component of the spectrum. The power spectral density may then be calculated using a fast fourier transform.

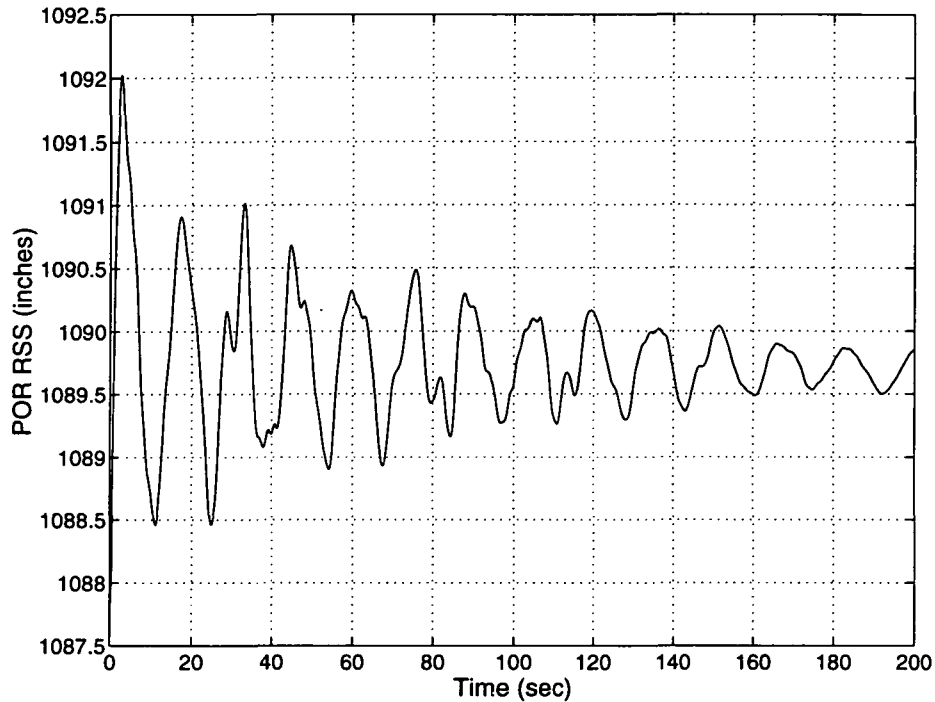


Figure 4.3: POR RSS time history for simulation output from 0.48 second -Roll, -Pitch PRCS jet firing.

The process does not stop there, however. The quality of the PSD that can be generated from a time history such as the one shown in Figure 4.3 is limited by both the sampling size and the total number of samples collected. The spacing between spectral lines in a PSD is called the resolution bandwidth, and is given by the formula

$$f_b = \frac{1}{N\Delta t}$$

where N is the number of samples and Δt is the spacing between samples. [14] So depending on the length of the time history being analyzed, we may have to

take further steps to ensure that the frequency resolution of the PSD is adequate to our purposes. The usual procedure in this sort of situation is to copy the data set and concatenate it onto the end of the original some number of times until the total length of the set is long enough to achieve the desired level of resolution. Of course, doing this inevitably introduces error into the spectrum, but over the years a variety of techniques have been developed to minimize this problem. In this particular case, the data set was copied, a hanning window was applied, and the data was concatenated five times to produce the graph shown in Figure 4.2 (a).

The above procedure was followed to generate the PSD's used for frequency identification and sequence selection. However, the number of concatenations varied from simulation to simulation, as well as the amount of data which had to be tossed out to remove the transients due to brake slippage. In general, PSD's were taken from simulations of short single pulses from the appropriate jets.

4.4 Simulation Results

DRS simulations were conducted using three different payloads representing the range of masses that might be encountered in current shuttle operations and in the near future. Simulations were also conducted with an unloaded arm in order to investigate possible applications of command preshaping to RMS-based photography or assistance of astronauts during EVA. Each of these payloads has its own unique features and has been used to

illustrate different aspects of constant amplitude command preshaping techniques.

4.4.1 The Upper Atmospheric Research Satellite

The Upper Atmospheric Research Satellite (UARS), pictured in Figure 4.4 extended on the shuttle's arm in the capture/release configuration, was deployed from the shuttle as part of STS-48 in September of 1991. A small to mid-sized payload, UARS weighed in at a little over 14,000 pounds. However, it was, in the words of the Draper engineers charged with pre-flight dynamic interactions analysis, "perhaps the most dynamic payload to be deployed from the Space Shuttle using the Remote Manipulator System." Its single large solar array was extremely flexible and possessed modes which could be easily excited by orbiter maneuvers, resulting in unacceptably high loads. A great deal of effort over a period of almost two years prior to the actual shuttle flight went into searching for a configuration of the shuttle autopilot that could be safely used with the satellite out on the end of the arm. A suitable application of command preshaping techniques could have made this task a lot easier.

The first step in implementing a constant amplitude preshaper is determination of the frequencies which must be removed. To this end, a DRS simulation was run to determine the response to a short (0.48 second) firing of three primary jets (F3U R2D R3D in Figure 2.1) producing an acceleration in the -Roll, -Pitch direction in the orbiter body axis system. The POR time histories were then taken from the simulation data, the root sum square calculated, and a PSD generated according to the procedure outlined in the

section above. The time history from this simulation can be seen in Figure 4.3, and the PSD in figure 4.2 (a).



Figure 4.4: The Upper Atmospheric Research Satellite deployed on the RMS in capture/release configuration.

Before proceeding on to analyze the simulation results, a few comments are in order. Throughout this results section PRCS simulations are based on -Roll, -Pitch jet firings and VRCS simulations are based on the firing of a jet which produces mainly +Yaw and +Pitch accelerations. These directions were chosen arbitrarily and kept consistent throughout for comparison purposes. They are not intended to represent "worst case" jet firings; what we are trying to show is the difference between the response to jets fired without preshaping and

jets fired with it. Though time was not available to show this conclusively, the improvement should be consistent in scale no matter which jets are fired.

Given the PSD in Figure 4.2 (a), it was decided to form a preshaping sequence centered on a frequency of 0.065 Hz. The preshaping sequence selected was an LY sequence, consisting of six pulses located at 0, $T/3$, $T/2$, $2T/3$, $5T/6$, and $7T/6$. The insensitivity curve for this sequence is shown in Figure 3.17. The LY sequence has several advantages which recommend it for this application. First of all, it has good frequency robustness around the design frequency, so it should do an excellent job of catching the two peaks visible in the PSD at 0.06 and 0.07 Hz. Secondly, it requires no double-size pulses, so the sequence can be implemented exactly in the constant amplitude framework, not approximated. Finally, the LY sequence is also a load sequence as described in Section 3.6. Minimization of loading in the arm should help to reduce brake slippage, and any reduction in the amplitude of the oscillation at the grapple fixture should also help to keep the loads down in the payload's solar array

The results of the application of an LY preshaping sequence to a UARS PRCS maneuver are summarized in Figure 4.5. The plot shows the time history of the POR X coordinate (the X direction showed the most excitation) for a total jet firing time of 0.48 seconds. The dotted line shows the result of firing the jets in a single pulse; the dashed line shows the result of firing the jets using the alt-mode delay of 3.44 seconds which was actually used in STS-48 [5]; and the solid line shows the result of firing the jets in an LY sequence. The LY sequence shows an average of a factor of 30 reduction in residual vibration over the single pulse, and a factor of 10 improvement over the alt-mode. This

is a clear vindication of the applicability of constant amplitude preshaping techniques to even complex nonlinear systems.

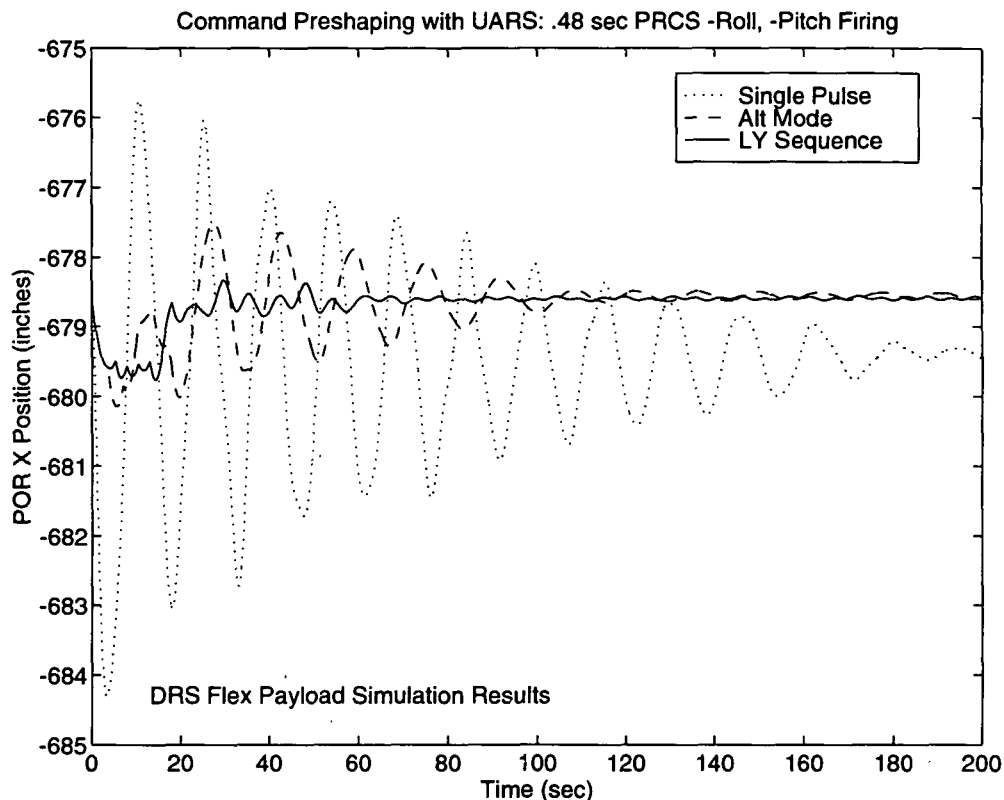


Figure 4.5: Simulation results from primary jet firings of total duration equal to 0.48 seconds. The dotted line shows the consequences of carrying out the entire command in a single pulse; the dashed line shows the results of using the alt-mode delay time between minimum impulse firings, and the solid line shows what happens when an LY sequence is used.

We have shown that an LY sequence can be used to reduce residual vibration given a good estimate of a system's fundamental frequency. We would also like to be able to show that the LY sequence offers substantial robustness to frequency error. With that end in mind, simulations were run where LY sequences were constructed around frequencies higher and lower than the frequency indicated by the PSD. One sequence was constructed about a frequency of 0.052 Hz, and another about a frequency of 0.081 Hz. The POR

time histories from these simulations are shown in Figure 4.6 compared with the time history from the sequence constructed about the best guess. Clearly, the LY sequence has extremely good robustness to frequency error. The curve for $\omega_N/\omega_D = 1.25$ actually does better than the nominal case in terms of residual vibration, and that deserves some comment. Power spectral densities do not identify a single frequency at which the system oscillates; they describe the distribution of energy in the system over a range of frequencies. The spreading of this energy over a range of frequencies can lead to unexpected results, and is one of the reasons that we must design sequences which eliminate vibration robustly.

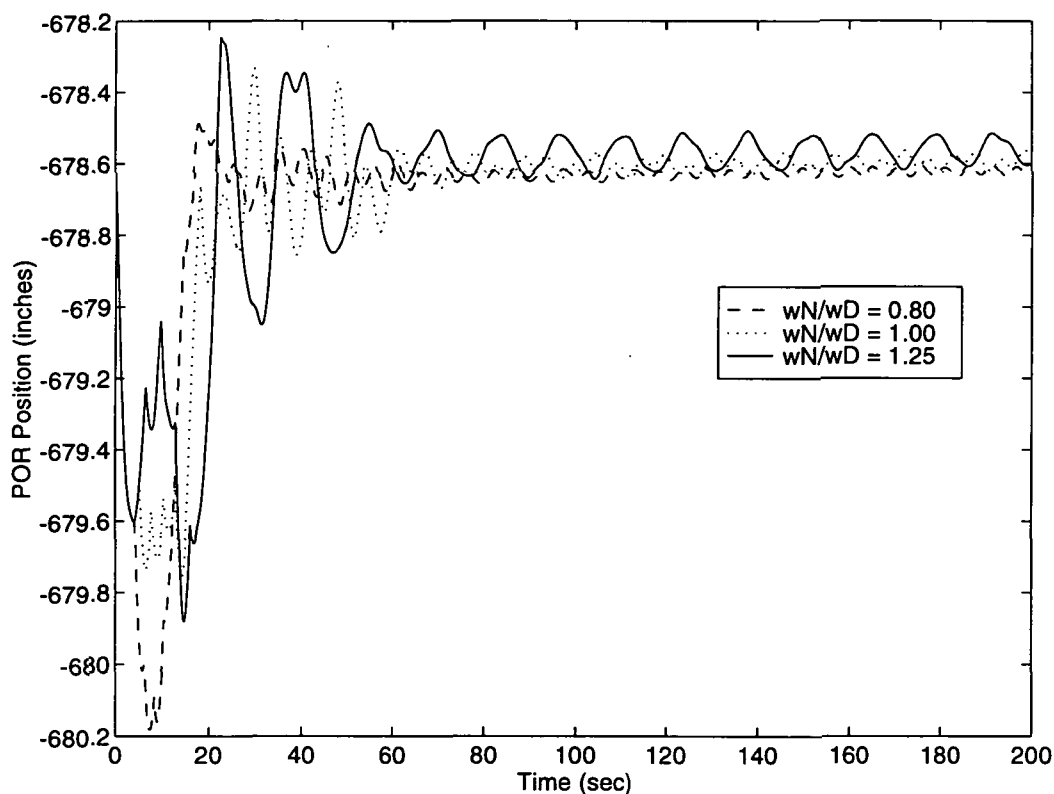


Figure 4.6: A test of LY sequence robustness. The solid and dashed lines show the time histories from sequences constructed for frequencies 20% lower and 25% higher than the frequency indicated by the PSD, respectively.

The above test of the LY sequence also provides an excellent opportunity to show that the principles behind the development of the load sequences described in Section 3.6 really work. The POR position should exhibit the same type of response as the loading: a disturbance from equilibrium with each jet firing followed by oscillation about the initial position. Figure 4.7 shows a close-up view of the time history plotted in Figure 4.6. The contrast between the response to the LY sequence and the alt-mode firings is quite distinct. The LY sequence keeps the deflection at a close to constant level throughout the jet-firing period. The alt-mode firings (see Chapter 2 for an explanation of the alt-mode), however, are spaced such that the deflection caused by one firing is sometimes aggravated by the next. The result is that peak deflections for the alt-mode are over 40% higher than for the LY sequence. Looking back to Figure 4.6, we can also see that the predictions of loads robustness made in Chapter 3 also have some merit. For $\omega_N/\omega_D = 1.25$, the peak displacement is only marginally higher than for the nominal case. As predicted in Section 3.6, this robustness only extends to one side of the nominal frequency; for $\omega_N/\omega_D = 0.8$ the peak is closer to that of the alt-mode.

Though the results above validate the principle behind the concept of load sequences, the practice is somewhat more complex. Figure 4.7 indicates that the LY sequence is minimizing loads in the arm, but it tells us nothing about the loads in the most fragile component of the system, the payload's solar array. Normally information about a payload's internal loading is not available from a DRS simulation; the simulation results are sent to another organization which uses them to drive a finite element model of the payload to determine the loading. In the case of UARS, however, Draper was able to

obtain a modal matrix which can be used to extract loading information from DRS flex payload output. Using this extra information, several simulations were run to determine whether command preshaping techniques could be used more effectively than the alt-mode to keep solar array loads down.

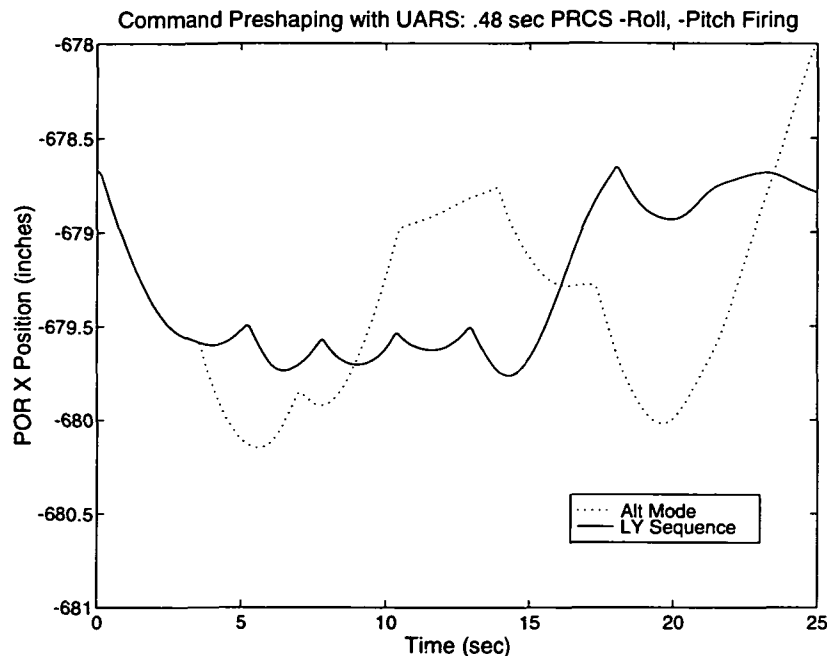


Figure 4.7: Close-up of the time history shown in Figure 4.5. The jet firings are spaced so as to keep the excitation in the arm to a minimum while the necessary impulse is being imparted.

The challenge here was to see if we could design a sequence short enough to be used with the high-impulse primary jets that could keep loads down at the solar array's frequency while at the same time keeping residual vibration in the arm down. This was partially successful. It was found that the standard LY sequence designed for the arm's fundamental frequency could do substantially better than the alt-mode at keeping the residual vibration down, but did not perform as well as the alt-mode in terms of minimizing solar array loads. An LY sequence designed for the solar array's fundamental frequency

gave a maximum solar array loading 33% lower than the alt-mode, but fared worse in terms of residual arm vibration. This sequence could be executed in one-fourth the time of an equivalent impulse alt-mode command. A compromise sequence designed to catch both frequencies produced about equal solar array loads and slightly better residual vibration, but took only half as long to execute as an equivalent impulse alt-mode command. In this type of situation, where two important modes are in operation, the ability to tailor sequences to specific needs in this manner and the generally lower time requirements (which translate into greater control authority) are the two main benefits which we can derive from command preshaping techniques relative to the alt-mode.

In general, the orbiter's low thrust vernier jets do not present as many problems as the primaries do. In the particular case of the UARS, loads in the solar panel due to VRCS jet firings were judged to be within an acceptable range. Still, it is informative to look at the performance improvement which could be achieved by using command preshaping techniques with the vernier jets. A simulation of a brief (0.96 sec) firing of jet F5L (see Figure 2.1) verifies that the fundamental frequency of vibration is almost the same as for the PRCS maneuvers described above. Figure 4.8 shows the simulation results for larger maneuvers of 9.6 seconds total firing time. The LY sequence was constructed based on a frequency of 0.07 Hz, as dictated by examination of the PSD of the test VRCS firing. Another simulation was run using an LY sequence based on the 0.065 Hz frequency used for the PRCS maneuvers. It produced almost exactly the same results, confirming the insensitivity of the sequence's performance to the small changes in fundamental frequency caused by

differing command directions. In any case, the LY sequence shows about a factor of 10 reduction in residual vibration over the single pulse firing. While not as spectacular an improvement as was observed for the PRCS firings, this is still quite good. One possible explanation for why it did not do better would be that the residual amplitude has become small enough to be dominated by normally insignificant nonlinear effects and disturbance torques. It may be impossible to do any better in this complex system.

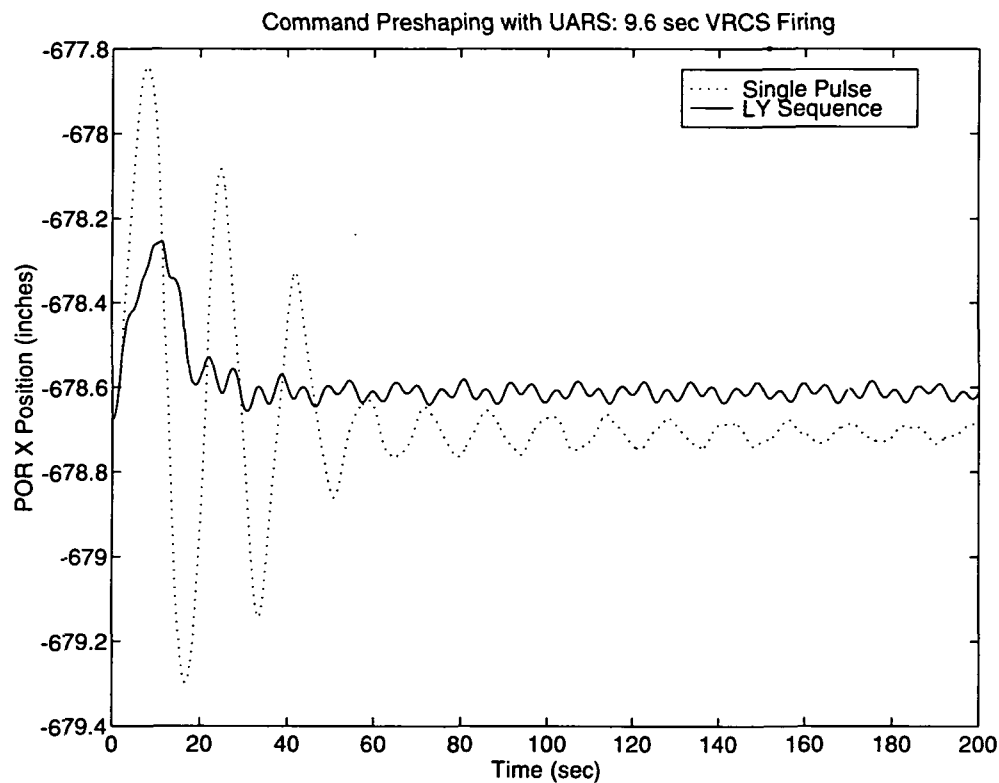


Figure 4.8: Time histories for VRCS jet firings totaling 9.6 seconds. The dotted line shows the response to a single pulse jet firing, while the solid line shows the response to firings conducted in an LY sequence.

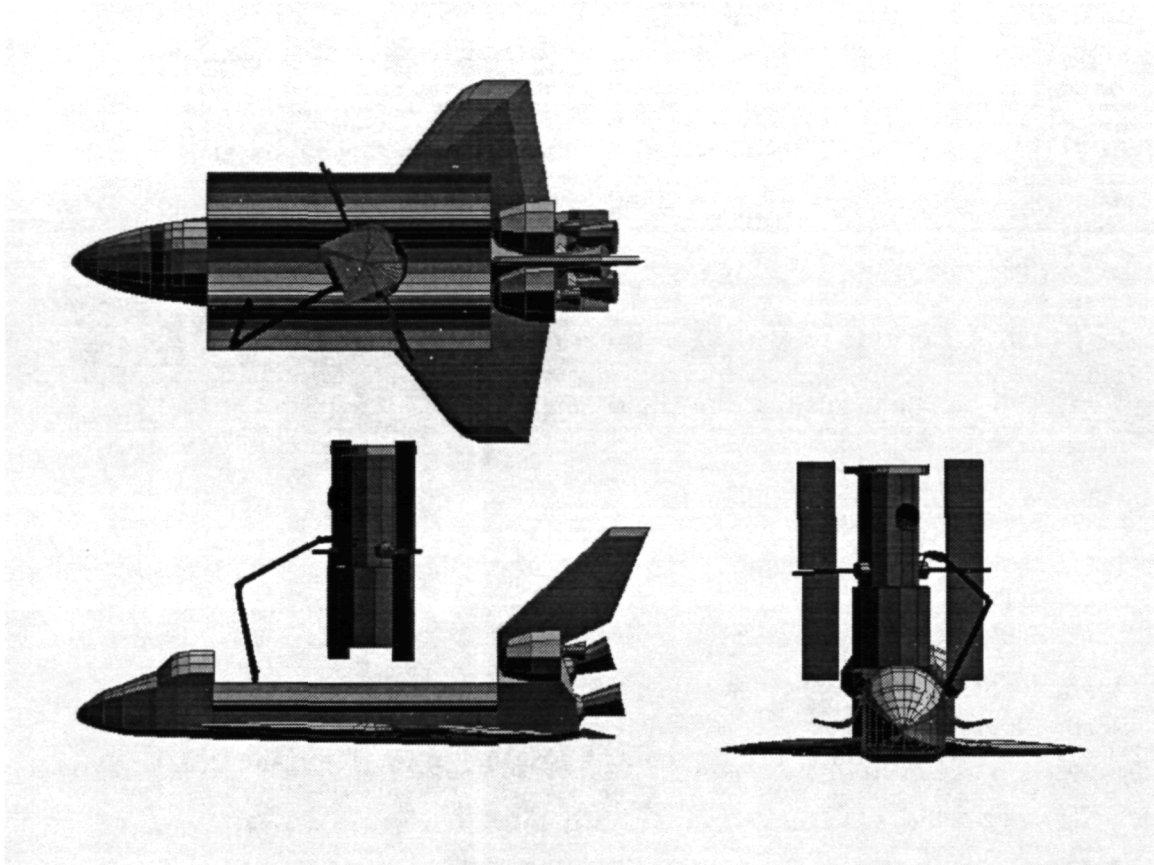


Figure 4.9: The Hubble Space Telescope deployed on the RMS in extended park position.

4.4.2 The Hubble Space Telescope

Just by itself, the Hubble Space Telescope (HST) makes a convincing argument as to why command preshaping techniques should be integrated into the shuttle autopilot. Shown in Figure 4.9 deployed on the arm in the extended park configuration, the 25,000 pound HST is representative of a class of heavy payloads which give NASA flight planners headaches. The large mass of payloads like HST exacerbates vibration and loading problems and makes dangerous RMS brake slippage much more likely. In addition, HST packs a pair of fragile solar arrays; during the initial EVA's in STS-61 (the

recent servicing mission) it was found that one of the old arrays had been bent to the point where it had to be jettisoned. Further incentive to implement some sort of improved control system is provided when we consider that HST has been out on the end of the arm in two missions already, and that NASA plans several more servicing missions over the telescopes operational lifetime. But the best argument for why an alternative control system is needed is provided by the record of the pre-flight analysis.

Despite the fact that considerable analysis had already been carried out for the space telescope in preparation for its initial launch on STS-31, dynamic interactions analysis for the recent servicing mission began at Draper a full six months before the STS-61 servicing mission. Considerable effort was expended over that time in searching for autopilot configurations which would allow maneuvers to be conducted without exciting too much vibration in the arm and violating the load limits on the telescope's solar arrays. That effort was ultimately only partially successful. A Draper memo released prior to the mission detailing recommended DAP configurations during RMS attached operations shows that no acceptable settings were found for PRCS operations (even under the alt-mode) with RMS brakes on, and that excessive loads were found in two of the four cases studied for VRCS operations with RMS brakes on. [19] Quoting from another Draper memo,

Potential violation of allowable load limits were indicated for three configurations, two for VRCS and one for alt-mode. All three violations were marginal and for off-nominal operations. For the alt-mode, an event causing a loads violation is extremely unlikely to occur even once during the flight. For the VRCS, an event producing similar jet pulsing to that used for the tests is possible, but unlikely. Based on engineering

judgment, it is very unlikely to occur more than once over an 8 hour period. [20]

In such a crucial mission to NASA's future, even a small possibility of such an occurrence is worrisome. Command preshaping techniques could be applied to relieve such fears and allow missions such as STS-61 to be carried out with less risk to the payload.

As always, the first step in applying constant amplitude preshaping is to find the frequencies at which the system vibrates. A DRS simulation was run to obtain POR time histories in response to a single minimum impulse firing of PRCS jets in the -Roll, -Pitch direction with the space telescope attached to the RMS in the extended park position pictured in Figure 4.9. The resulting data was used to generate the PSD shown in Figure 4.10. A fundamental frequency of 0.06 Hz is clearly visible in the power spectrum, as is a secondary frequency at about 0.115 Hz.

Given the shape of the power spectrum, the LY sequence again seems to be an ideal candidate for testing of preshaping using PRCS jets. As can be seen from Figure 3.17, the LY sequence has a second valley at twice the design frequency, so it should go some way toward removing the second frequency in the spectrum shown above. The simulation results from a total firing time of 0.54 seconds divided up in three different ways are shown in Figure 4.11. The alt-mode delay time used was the value of 12.1 seconds (roughly $3T/4$) recommended for the actual HST servicing mission.

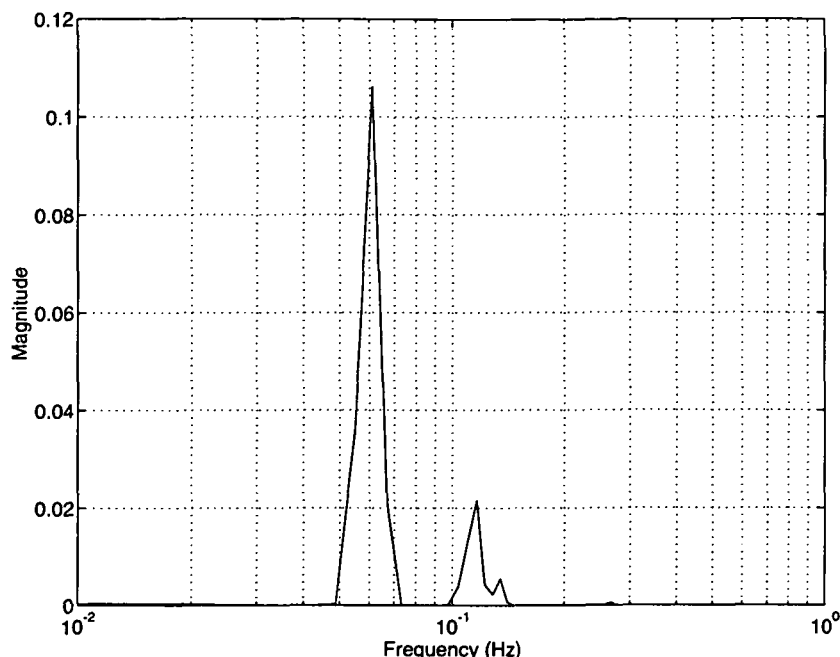


Figure 4.10: Power Spectral Density of POR response to a 0.08 second -Pitch, -Roll PRCS jet firing with HST deployed on the arm in the extended park configuration.

The most noticeable aspect of the simulation results is, of course, the dramatic difference between the mean POR values of the single pulse response and the response to the other two firing patterns. This is due to brake slippage in the RMS; when the arm is subject to sudden accelerations the brakes are designed to slip in order to relieve the loads and prevent damage to the arm. Both the LY sequence and the alt-mode fire jets in a more spread out fashion, so accelerations are more gradual and less brake slippage occurs.

In terms of residual vibration, the LY sequence clearly does the best. The performance increase is not as stellar as it was for UARS, however. One likely explanation for this is that the brake slippage in the single pulse run acts to reduce the vibration in the arm as well as the loads because it removes energy from the system. Simulations run with artificially high brake friction to

prevent brake slippage exhibited substantially larger oscillations than those where the brakes were allowed to slip. Thus, in the simulation runs with standard brake friction the LY sequence is being compared to a vibration level that has already been reduced by brake slippage.

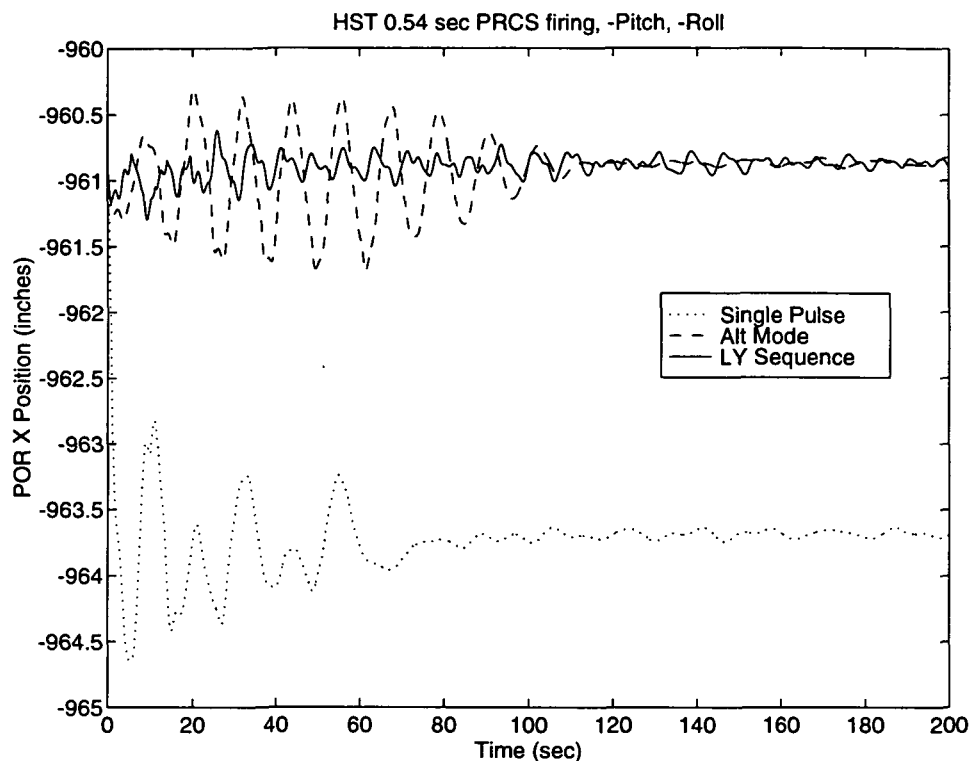


Figure 4.11: Time histories for PRCS jet firings totaling 0.54 seconds. The dotted line shows the response to a single pulse; brake slippage of over 3 inches is visible. The dashed line shows the response to jet firings distributed according to the alt-mode's delay time, and the solid line shows the response to pulses distributed in an LY sequence.

The important comparison to make here is between the preshaping sequence and the alt-mode. In this case the LY sequence provides about a factor of two reduction in residual amplitude over the alt-mode, but it does so with a substantially shorter command. Where the alt-mode takes 60 seconds to complete the command, the LY sequence takes a little under 20 seconds. This

reduction in execution time translates into substantially greater control authority for the shuttle.

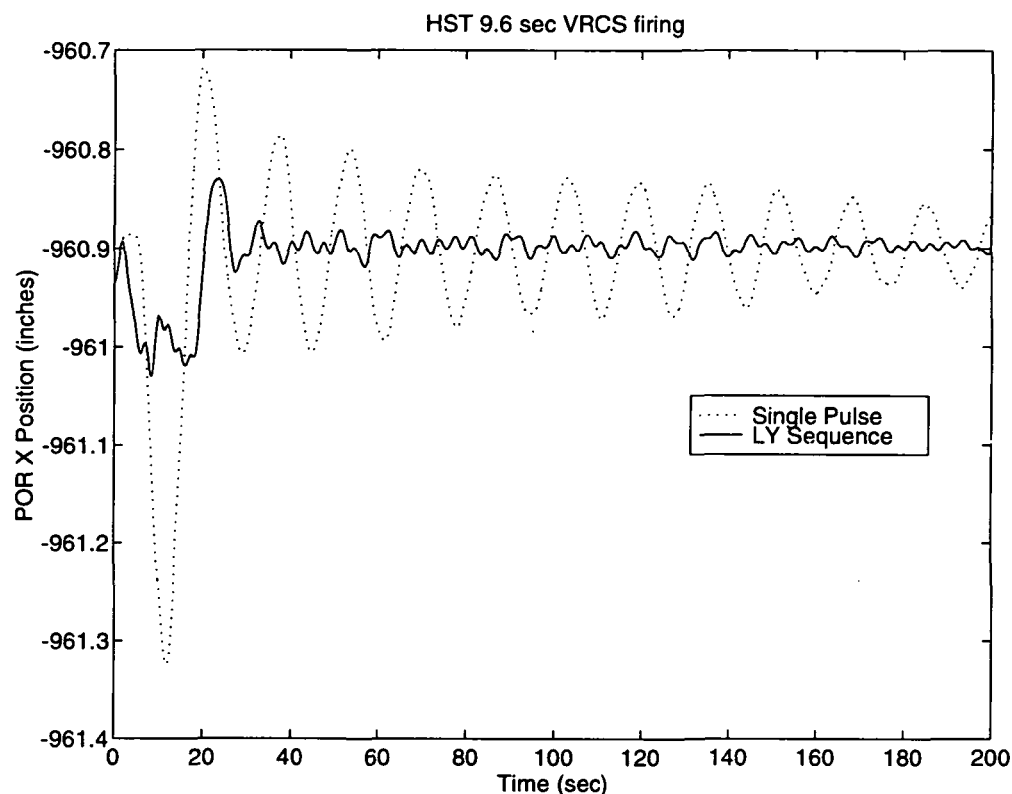


Figure 4.12: Time histories for VRCS jet firings totaling 9.6 seconds. The dotted line shows the response to a single pulse. The solid line shows the response to pulses distributed in an LY sequence.

Shifting from an analysis of PRCS operations to a look at the vernier jets, we find that the LY sequence is still an attractive starting point. After confirming that the frequencies of concern were unchanged by the switch in jets, two simulations were run to compare the POR response to an LY sequence with the response from a single pulse. The jet fired was F5L (see Figure 2.1), and the total burn time was 9.6 seconds. The resultant POR X coordinate time histories are shown in Figure 4.12. The graphs show that the LY sequence

provides a substantial reduction in residual vibration, on the order of a factor of 10.

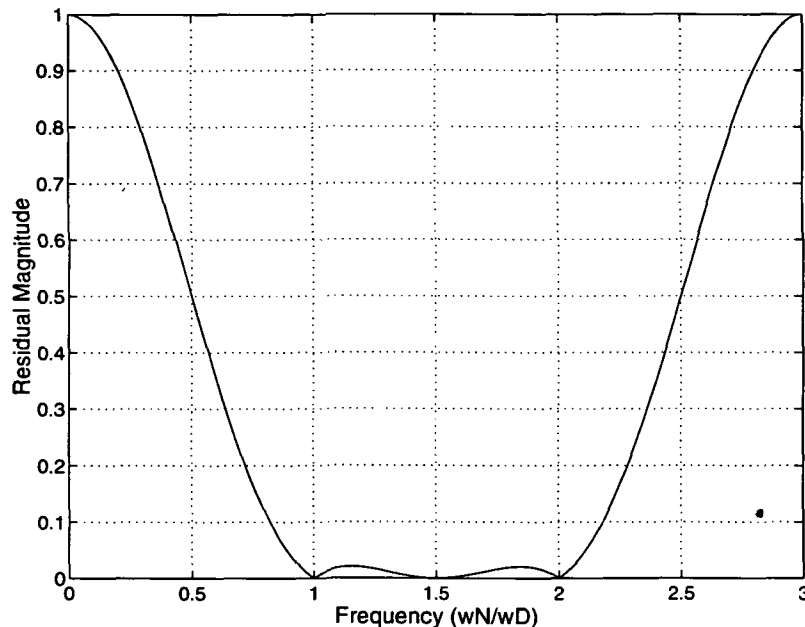


Figure 4.13: The insensitivity curve of an $L^2@3/2Y$ sequence.

At this point, however, we can safely say that the LY sequence suffers from one serious flaw from the point of view of this thesis: it's boring. Its effectiveness has been well established, so it is time to test out some other sequences. Since we have a fair amount of room to experiment with the longer VRCS firings, it seems like a good opportunity to try out one of the more complex sequences. For starters, we shall try out a sequence which convolves a Y sequence with an L^2 sequence designed for a slightly higher frequency. Since the L^2 sequence is designed for a frequency 1.5 times higher than the Y sequence, we shall label the full sequence $L^2@3/2Y$, an informative designation, if not a pretty one. The sequence's insensitivity curve is shown in Figure 4.13. The sequence attenuates anything between the design frequency and twice its value, and thus should catch that second peak in the

PSD shown in Figure 4.10, which comes in at a frequency just a bit under twice that of the fundamental mode at 0.06 Hz.

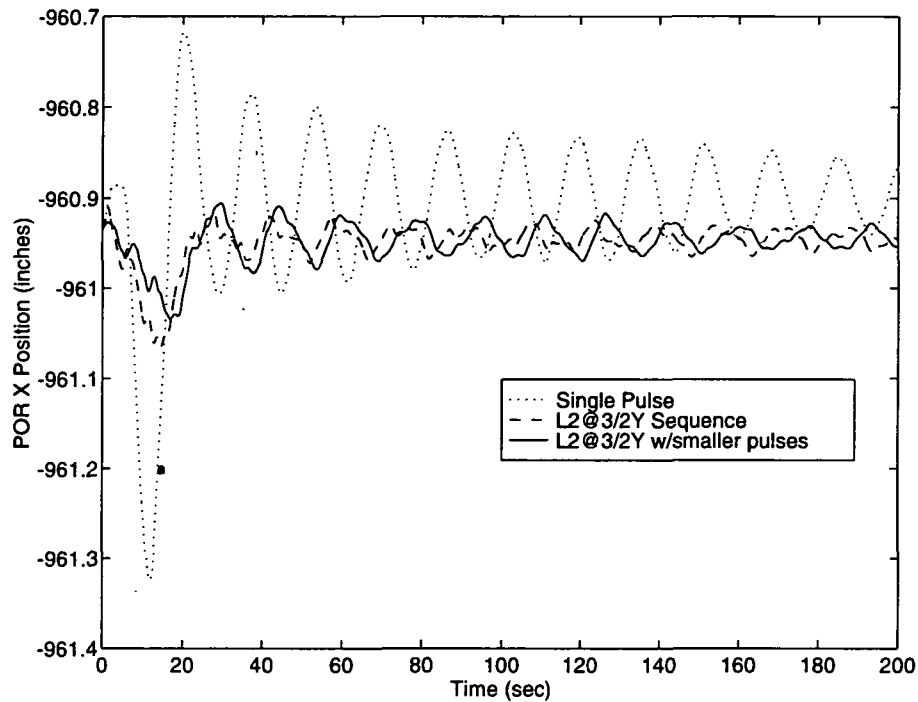


Figure 4.14: A look at the effect of pulse size on preshaping sequence performance. The solid line shows the performance of an $L^2@3/2Y$ sequence using smaller width sequences concatenated together. The dashed line shows the performance of the sequence using full width pulses. The dotted line showing the performance of a single pulse VRCS jet firing is shown for reference purposes.

The impulses of the $L^2@3/2Y$ sequence come in a 1-3-4-3-1 pattern, with a separation of a third of period between each impulse. The varying impulse sizes mean that the larger impulses will have to be approximated with wider pulses. This gives us an excellent opportunity to test whether the methods discussed in Chapter 3 for keeping pulse size down have any effect in a real system. Given a 9.6 second total firing, the amplitude 4 impulse by default will be approximated by a pulse 3.2 seconds long, which is about 20% of the length

of the system's period. This should result in a noticeable degradation in performance. For comparison purposes, another simulation was run using a 5 times smaller pulse width, with the smaller sequences concatenated together as in Figure 3.9 to provide the same total firing time. The results of the two simulations are shown in Figure 4.14, compared with the response to a single pulse of equivalent firing time. Clearly, benefits of a smaller pulse size have been obscured by the system's innate complexities. In comparison with the LY sequence, we can see that the $L^2@3/2Y$ sequence has a little higher average residual amplitude but less brake slip and no overshoot at the end of the sequence such as can be seen in Figure 4.12.

4.4.3 Space Station Model MB6

The space shuttle will play a key role in the construction of the planned international space station over the next decade. The RMS will be used to help with each phase of construction, and plans to use the RMS to aid in docking have also been mooted. With such a large payload on the end of the arm there are bound to be severe difficulties with vibration, especially when one considers the extremely flexible nature of the planned station and the low frequencies of its modes. That the original design specifications for the RMS call for a maximum payload weight of 65,000 lbs does not help either, especially when the mass of the station *without* most of its modules is more than twice that limit. RMS attached space station operations represent a key opportunity for the application of command preshaping techniques.

The primary difficulty with assessing the effectiveness of any control technique in conjunction with the space station is, of course, that the station keeps changing. For the purposes of this thesis, a space station model from late 1992 has been used, code-named MB6. The MB6 station configuration is pictured in Figure 4.15. It represents the configuration of the station at the start of the sixth station construction mission, right before it enters its maintained operational mode. The power system, main truss, and docking port have been assembled and the lab module is about to be added on.

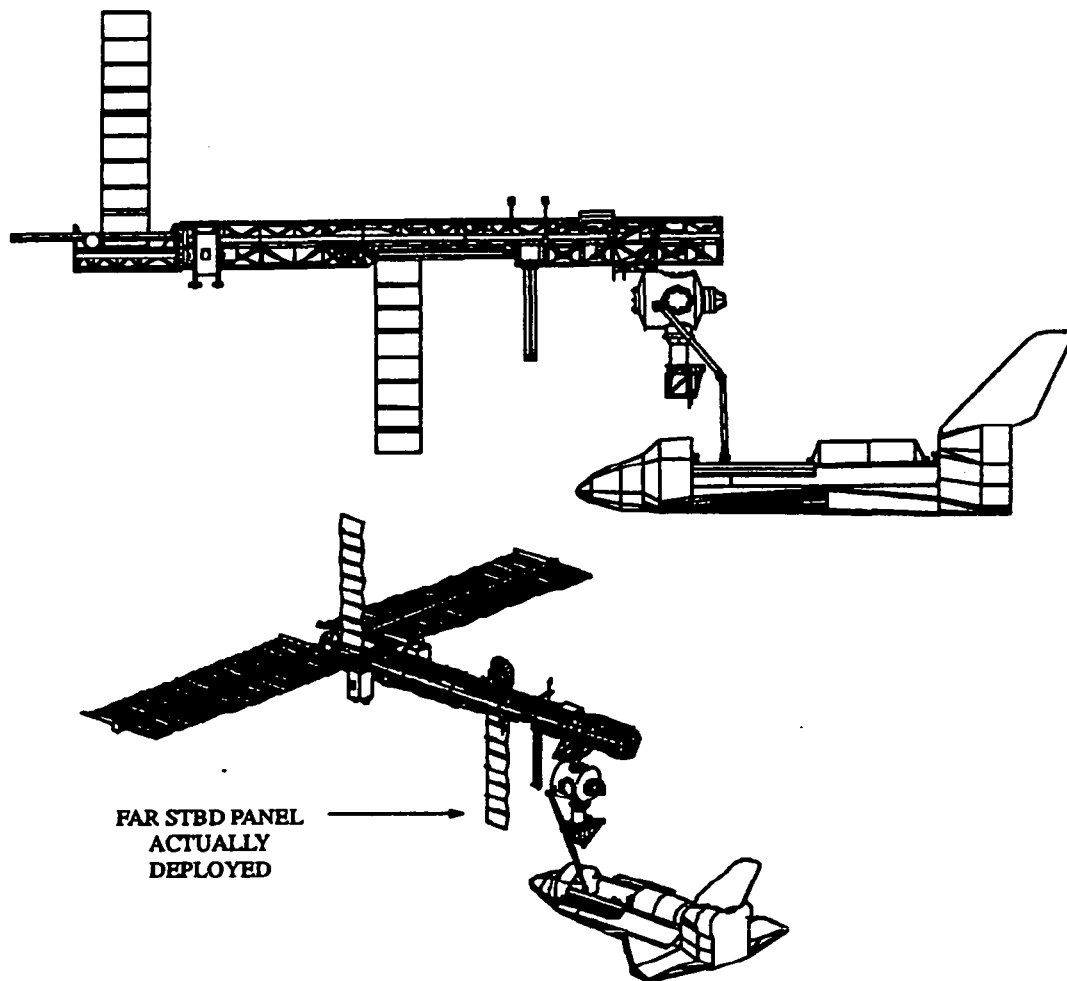


Figure 4.15: The MB6 space station configuration with shuttle attached via the RMS as planned in 1992.

The large mass of the station, in whatever configuration is finally flown, will combine with its flexible characteristics to make it extremely unmanageable when attached to the orbiter via the RMS. Figure 4.16 (a) shows a PSD of the POR response to a short VRCS firing. Note that there are four distinct frequencies visible: two at 0.009 and 0.011 Hertz and two more at 0.05 and 0.07 Hz. With the alt-mode, which can only handle a single frequency, choosing an appropriate delay time would be extremely difficult, if not impossible. With the constant amplitude preshaping techniques presented in the previous chapter, designing a firing sequence which catches all of these frequencies is quite simple. We simply take the sequence shown in Figure 3.16 centered about 0.06 Hz, convolve it with an L^2 sequence designed about a frequency of 0.01 Hz, and we arrive at a 12 pulse sequence with an insensitivity curve shown in Figure 4.16 (b). The insensitivity curve is juxtaposed with the power spectral density to show how each of the important frequencies is removed.

The 12 pulse sequence described above is probably too large to be applied efficiently to the PRCS jets. The large number of pulses combined with the high impulse of the primary jets would likely result in unacceptable command resolution. Instead, it was tested with a total VRCS firing time of 38.4 seconds — long enough to get up to a reasonable maneuver rate even when the large mass of the station is considered. The results of the simulation are shown in Figure 4.17. Despite the presence of modes at four different frequencies, we still manage to get a factor of 10 improvement. This improvement is especially nice since the oscillations from the single pulse firing show little sign of

damping. Without the use of preshaping techniques, significant vibration could easily persist for half an hour after the maneuver was completed!

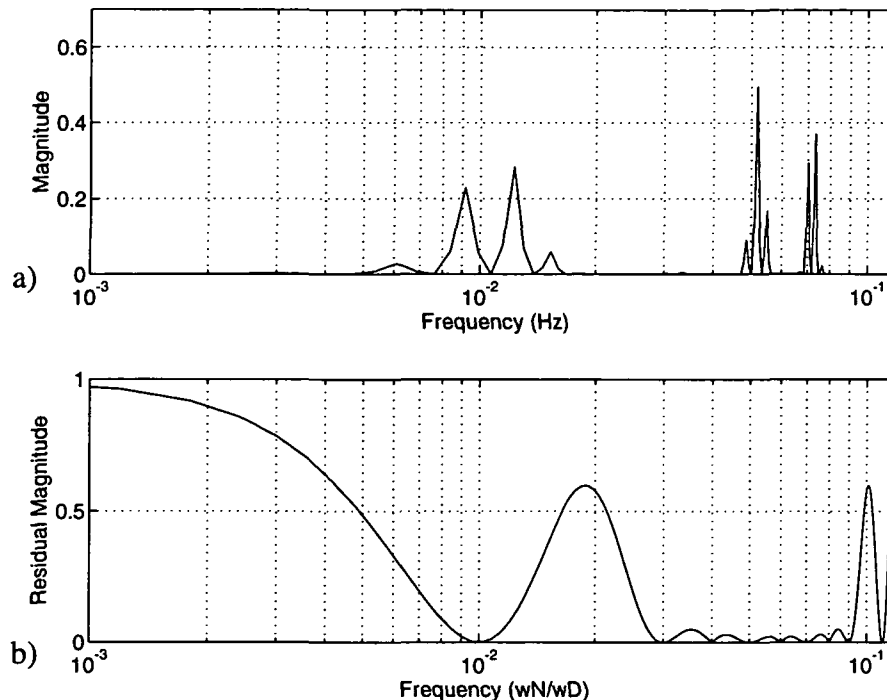


Figure 4.16: a) Power Spectral Density of MB6 POR time history from simulation of a short VRCS firing. b) Insensitivity curve of a sequence designed to attenuate all four modes visible in (a).

Though we cannot use sequences as complex as the one described above for PRCS operations, we can still achieve substantial improvements. Two simulations were run using the MB6 model with a 0.64 second total firing time. One simulation executed the entire maneuver with a single pulse firing, while the other divided the 8 pulses up using an L^2 sequence centered at 0.01 Hz convolved with an L sequence centered at 0.05 Hz. POR time histories from the two simulations are shown in Figure 4.18. No comparison was made to alt-mode performance, as no appropriate alt-mode delay time could be found. Whereas the single pulse firing causes the brakes to slip enough to cause a shift of over 2 feet in POR position, the preshaped firing results in no significant slippage

at all. The improvement offered in residual amplitude by the preshaped sequence is relatively minor, but then again most of the energy imparted to the arm from the single firing was taken out by brake slippage, so the preshaped sequence is really doing quite well. It must be emphasized that here command preshaping techniques are making possible operations which probably could not be conducted with the current alt-mode upgrade.

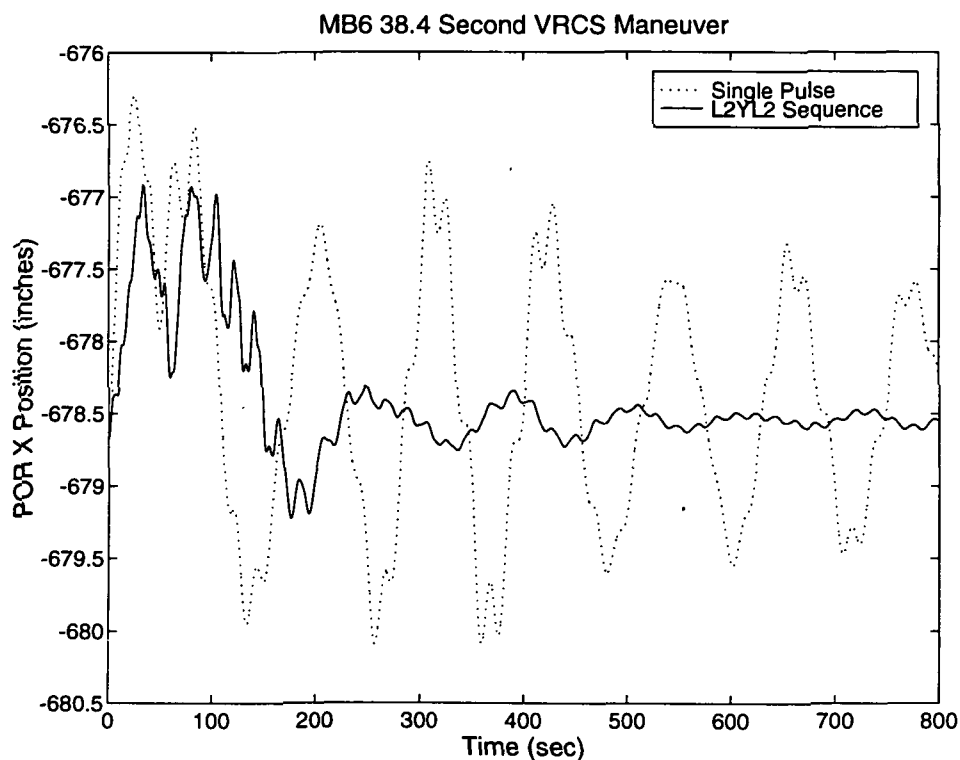


Figure 4.17: Time histories for VRCS jet firings totaling 38.4 seconds. The dotted line shows the response to a single pulse. The solid line shows the response to pulses distributed in the sequence with insensitivity curve shown in Figure 4.16 (b).

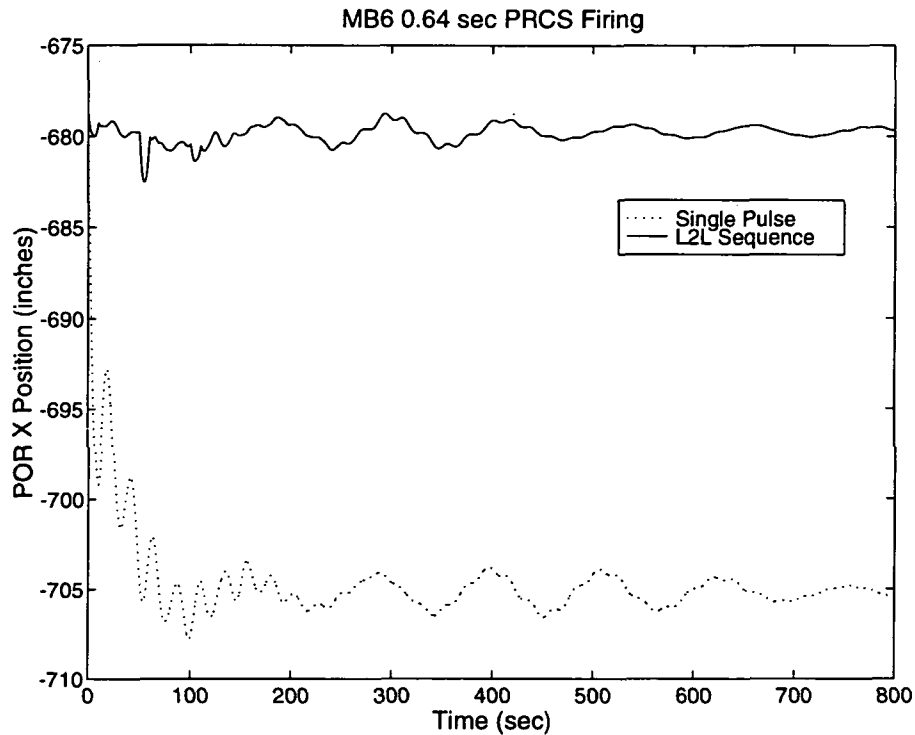


Figure 4.18: Time histories for PRCS jet firings totaling 0.64 seconds. The dotted line shows the response to a single pulse; note that brake slippage causes a shift in mean POR position of over 4 feet. The solid line shows the response to pulses distributed in an LY sequence.

4.4.4 Unloaded Arm

As has been mentioned in previous sections, there may be occasions where vibration in an unloaded arm due to RCS jet firings is a problem. This could occur when an astronaut is using the RMS as a platform for some sort of EVA work (the astronaut's mass is considered to be negligible), or perhaps when a stable camera shot is desired for a NASA video. A quick test shows that the techniques demonstrated above for larger payloads apply equally well to an unloaded arm. A DRS simulation of a single minimum impulse PRCS firing was used to generate a PSD of the arm's POR motion; this yielded a fundamental frequency of about 0.46 Hz. Centering an LY sequence about this frequency,

another simulation gave the results shown in Figure 4.19. Clearly, oscillations are damped out more rapidly when no payload is attached to the arm, but for the first 10 seconds or so of the response we can see a marked reduction in residual vibration. Differences on this sort of time scale could be quite significant for the type of operations where no payload is attached to the arm.

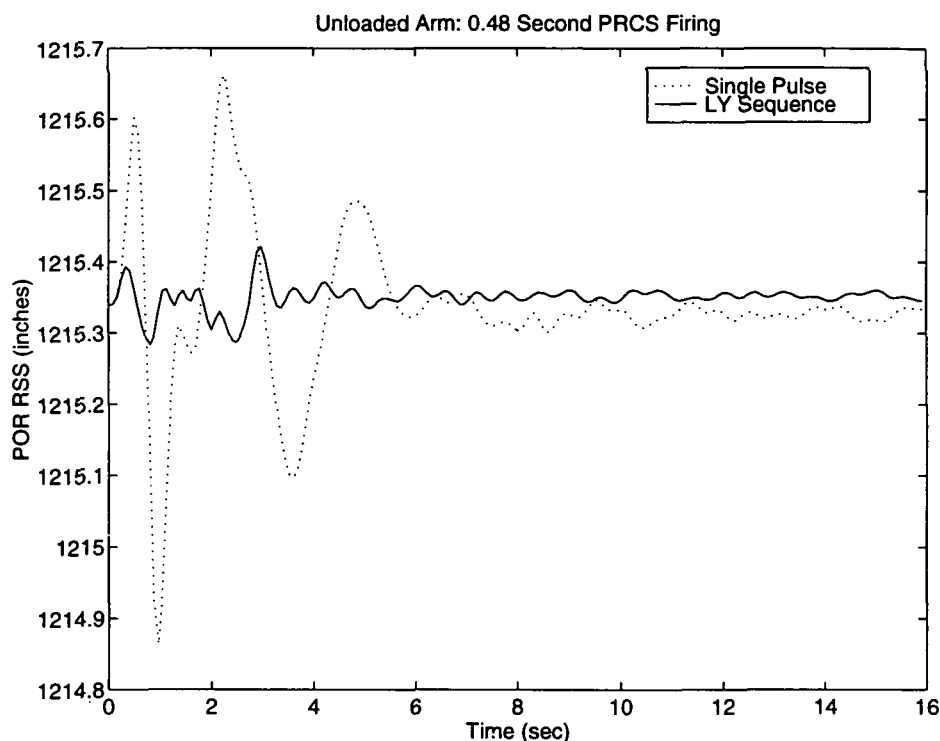


Figure 4.19: Time histories for PRCS jet firings totaling 0.48 seconds for an unloaded arm. The dotted line shows the response to a single pulse. The solid line shows the response to pulses distributed in an LY sequence.

Conclusion

Chapter 5

When work first started on this topic, it was unclear how severe the limitations of a constant amplitude control system would prove to be. One school of thought held that conversion from the continuous domain to the constant amplitude domain would be a simple matter; another was uncertain that any but the most basic concepts of command preshaping could be applied to a system with so few degrees of freedom. The results have, in general, been promising, and in many cases not quite so simple as one might have thought. In addition, the restrictions involved in working with constant amplitude actuators have resulted in some unexpected insights that can be applied to the continuous domain as well.

In Chapter 3, we examined the theory of constant amplitude preshaping. We started with the basic principles which determine the response of a flexible system to a single pulse, and then used that knowledge to adapt preshaping techniques developed for the continuous domain to serve in a

system with constant amplitude actuators. We discussed the compromises which must be made to approximate impulse sequences with varying amplitudes and the changes in viewpoint necessary to measure the performance of CAP sequences. New types of sequences were proposed which, at the cost of slightly greater time cost, could be used to achieve robustness to frequency error without having to use impulses of different magnitudes. These sequences also offered the possibility of attenuating multiple frequencies with a smaller time cost than that imposed by convolving together standard impulse sequences. Another type of sequence was described that, in addition to eliminating residual vibration, is capable of minimizing the loads experienced by a system during the time in which a command is carried out. Issues of damping and multiple mode systems were addressed, and a process presented by which preshaping sequences could be constructed to fit the constraints of individual problems.

Three aspects of the theory presented in this thesis are of particular importance and deserve special note. First of all, the approach taken to sequence design in this thesis was quite different from that taken in most of the previous work on the subject of command preshaping. That work has focused to a large extent on selection of sequences using nonlinear optimization techniques which provide good results but little insight; in this thesis, sequences were built up from basic principles and combined and positioned to achieve the desired effect. Admittedly, this is easier to do in a system with such limited degrees of freedom, but the insights provided by this method of sequence design can often be quite valuable. Secondly, the use of sequences with greater numbers of pulses but smaller spacing has great

potential. These sequences have tended to be overlooked in previous work because they incur larger time delays than sequences based on a half-period spacing, but in cases where system modes fall in whole number ratios they can offer equivalent robustness with smaller delays than sequences generated by convolution. Finally, the sequences developed for minimizing system loads show great promise. The ability to conduct a maneuver with minimal excitation during the move, no residual vibration when the move is over, and significant robustness to frequency error has a wide range of potential applications in a number of fields.

Application of the theory to a high fidelity model of a real system went extraordinarily well. Command preshaping techniques were applied to space shuttle attitude maneuvers while a variety of payloads, including the Upper Atmospheric Research Satellite, the Hubble Space Telescope, and one possible configuration of the proposed space station, were extended on the shuttle Remote Manipulator System. Simulation results showed reduction in vibration by about a factor of 30 when the preshaping sequences were compared with single pulse commands, and by as much as a factor of 10 when compared with the Alt-Mode. Comparison of that same simulation's output with flight data from the Space Telescope servicing mission showed virtually no discrepancy and added greatly to the credibility of the results. Further simulations demonstrated the robustness of the preshaping sequences to frequency error and their ability to reduce system loads. In one case where use of the alt-mode would be difficult or impossible, preshaping techniques were able to prevent brake slippage which would result in an end-point displacement of over two feet if a single pulse firing was used.

5.1 Future Work

The eternal dilemma of scientific research is that for each question that is answered, two more rise up to take its place. This work is no exception, and there are many questions which remain to be answered.

So far, most of the work conducted in the area of command preshaping has relied on nonlinear optimization for sequence selection. In this thesis, a heuristic approach was used instead. No serious comparison of the effectiveness of the two methods was made, however. It would be useful to see which performs best under a variety of circumstances and how a combination of the two methods compares with each alone.

Some investigation of the analytical aspects of command preshaping was conducted for this thesis. For several different sequences the shapes of the insensitivity curves were pinned down with specific equations and the effect of pulse width on approximation of non-unit magnitude impulses was codified. Much work remains to be done in this area, however. Precise equations showing how variation in specific sequence parameters affects performance will offer the designer far greater control over the system's behavior. Also, a closer investigation of the magnitude-phase relationship in the insensitivity curve may yield some interesting insights.

The concept of load-minimizing preshaping sequences was presented in this thesis but could use substantial further work. A formulation able to cope with multiple frequencies would be especially useful.

On a more practical note, a worthy project would be the development of software for sequence design. Automatic sequence selection for a given set of system specifications and interactive design capabilities should both be available. Such software should also allow the use of both nonlinear optimization techniques and heuristic methods.

Experimentally, there are two major tasks that it would be nice to see carried out. The first, and easiest, would be to verify the pulse width effects discussed in Chapter 3 on a simpler system than the space shuttle. The second, and much more difficult task, would be to integrate constant amplitude preshaping techniques into the software of the space shuttle's Digital AutoPilot.

Integrating CAP techniques into the DAP will be difficult because the DAP is a closed loop control system, and by current theory constant amplitude preshapers are fundamentally open-loop constructs. It is probably possible to get around this, but it will surely take substantial thought to achieve the seamless integration necessary to ensure safe operation under all conditions. As a first test, the easiest point of entry for preshaping techniques would appear to be in the DAP's manual discrete rate mode. After that has been proved to work successfully, a more complete integration into the autopilot could be attempted.

Bibliography

- [1] Adams, Neil, "Worst Case Directions for Exciting High End Effector Accelerations with the HST in the Capture/Release Configuration", CSDL Memo SSV-93-077, September 1993.
- [2] Appleby, Brent, *Reducing Shuttle-Payload Dynamic Interaction with Notch Filters*, Master of Science Thesis, Massachusetts Institute of Technology, CSDL-T-939, 1987
- [3] Barrington, Ray, "STS-61 FCS Flight Report", CSDL Memo SSV-94-002, January 1994.
- [4] Brown, Mark Everett, *Rapid Slewing Maneuvers of a Flexible Spacecraft Using On/Off Thrusters*, Master of Science Thesis, Massachusetts Institute of Technology, CSDL-T-825, 1983.
- [5] DeBitetto, Paul A., and David Hanson, *STS-48 UARS Optional Service Final Report: Solar Array Loads Analysis*, CSDL-R-2379, Draper Laboratory, 1992.
- [6] G&C Systems Training Manual: Guidance and Flight Control — Insertion, Onorbit, and Deorbit; I/O/D G&C 2102; NASA, 1985.
- [7] Gray, Carl, "STS-61 HST Flight to Simulation - Low Hover to Release Orbiter Motion", CSDL Memo SSV-94-001, January 1994.
- [8] Hain, Roger, "Stability Screening of Russian Alpha flight 2R", CSDL Memo SSA-93-10, November 1993.
- [9] Hanson, David P., "Comments Regarding STS-31 PDRS Crew Debriefing", CSDL Internal Memo, EGC-90-141, July 1990.
- [10] Henderson, Tim, "HST Servicing Mission Simulation Models (RMS Models)", CSDL Memo ESC-93-212, August, 1993.
- [11] Hyde, James, M., *Multiple Mode Vibration Suppression in Controlled Flexible Systems*, Massachusetts Institute of Technology Artificial Intelligence Laboratory Technical Report No. 1018, 1988.
- [12] Kirchwey, Kim, "STS-61 RMS-Attached Stability, Performance and Release Analyses", CSDL Memo SSV-93-095, January 1994.
- [13] Knapp, Roger, *Fuzzy Based Attitude Controller for Flexible Spacecraft with ON/OFF Thrusters*, Master of Science Thesis, Massachusetts Institute of Technology, CSDL-T-1183, 1993.

- [14] Mayhew, D., "The Estimation of Power Spectral Density from Sampled Data Records of a Stochastic Process", CSDL-C-5595, Draper Laboratory, 1983.
- [15] Metzinger, Richard W., *DRS Users Guide*, CSDL, 1993.
- [16] Palleson, Don, "STS-31 PDRS Post Mission Report", 1990.
- [17] Rappole, Bert W., Jr., *Minimizing Residual Vibrations in Flexible Systems*, Master of Science Thesis, Massachusetts Institute of Technology, CSDL-T-1144, 1992.
- [18] Sackett, Les, "Acceptability of Excessive Loads for RMS Operations on STS-61", CSDL Memo SSV-93-114, November 1993.
- [19] Sackett, Les, "STS-61 DAP Recommendation Summary", CSDL Memo SSV-93-110, November 1993.
- [20] Sackett, Les, and Dave Hanson, "Assessment of RMS-attached Forcing Function Conservatism for HST SM Solar Array Loads Analysis", CSDL Memo SSV-93-108, November 1993.
- [21] Shuttle Flight Operations Manual, Vol. 16 - Payload Deployment and Retrieval System, JSC-12770, NASA, 1981.
- [22] Shuttle Operational Data Book, NSTS-08934, Vol. 1, Revision E, 1988.
- [23] Singer, N. C., and W. P. Seering, "Preshaping Command Inputs to Reduce System Vibration", *Journal of Dynamic Systems, Measurement, and Control*, Vol. 112, March 1990, pp. 76-81.
- [24] Singer, Neil C., *Residual Vibration Reduction in Computer Controlled Machines*, Doctoral Thesis, Massachusetts Institute of Technology, AI-TR 1030, 1988.
- [25] Singhose, William, "Shaping Inputs to Reduce Vibration: A Vector Diagram Approach", Massachusetts Institute of Technology Artificial Intelligence Laboratory, Memo #1223, 1990.
- [26] Smith, O.J.M., *Feedback Control Systems*, New York: McGraw-Hill, 1958.
- [27] Stoddard, Ike, "Flight Cycle On-Orbit DAP I-Loads and Generic Stability Envelopes", memo for STS-61, ESC-93-188, August 6, 1993.
- [28] STS-37 Flight Plan, JSC-23619-37, NASA, December 19, 1990.
- [29] STS-61 Preliminary Flight Plan, JSC-48000-61, NASA, February 19, 1993.
- [30] Turnbull, Joseph F., and Phillip D. Hattis, *State Estimator Upgrades Certification Plan*, CSDL-R-2477, 1992.
- [31] Wie, Bong, and Ravi Sinha, "Robust Fuel- and Time-Optimal Control of Uncertain Flexible Space Structures", pp. 2475-2479, Proceedings of the American Control Conference, San Francisco, CA, 1993.

- [32] Zimpfer, Doug, "STS-31 Transition/On-orbit FCS Flight Performance Evaluation", CSDL Memo SSV-90-025, May 1990.
- [33] Zimpfer, Doug, and Ray Barrington, "STS-37 Transition/On-orbit FCS Flight Performance Quick Look Report", CSDL Memo SSV-91-018, April 1991.