

Optimal control design of band-selective excitation pulses that accommodate relaxation and RF inhomogeneity

Thomas E. Skinner^{a,*}, Naum I. Gershenson^a, Manoj Nimbalkar^b, Steffen J. Glaser^b

^a Physics Department, Wright State University, Dayton, OH 45435, USA

^b Department of Chemistry, Technische Universität München, Lichtenbergstr. 4, 85747 Garching, Germany

ARTICLE INFO

Article history:

Received 17 November 2011

Revised 11 February 2012

Available online 24 February 2012

Keywords:

Selective excitation

RC-SEBOB

Relaxation

T_1 relaxation

T_2 relaxation

Optimal control theory

ABSTRACT

Existing optimal control protocols for mitigating the effects of relaxation and/or RF inhomogeneity on broadband pulse performance are extended to the more difficult problem of designing robust, refocused, frequency selective excitation pulses. For the demanding case of T_1 and T_2 equal to the pulse length, anticipated signal losses can be significantly reduced while achieving nearly ideal frequency selectivity. Improvements in performance are the result of allowing residual unrefocused magnetization after applying relaxation-compensated selective excitation by optimized pulses (RC-SEBOBs). We demonstrate simple pulse sequence elements for eliminating this unwanted residual signal.

© 2012 Elsevier Inc. All rights reserved.

1. Introduction

Frequency-selective pulses have widespread utility in magnetic resonance and have motivated significant efforts towards their design [1–47]. In many useful cases, the resulting methodologies can achieve the best approximation to the ideal rectangular profile for spin response as a function of frequency offset.

However, in all of these approaches to pulse design, performance criteria that can be included in the design protocol are restricted either by analytical procedures of highly specific scope or by numerical methods that are limited by the efficiency of the optimizations employed. As a result, pulse response is typically optimized only for the nominal ideal RF pulse values. In addition, although the length of pulses required for narrowband applications can significantly reduce their effectiveness when relaxation times are comparable to the pulse length [48,49], the solution to the problem—selective pulses which are less sensitive to relaxation effects—can also be demanding for standard design methods [33,50–54].

To make these design challenges tractable, the space of possible pulse shapes is often reduced by forcing the solution to be a member of a particular family of functional forms (for example, finite Fourier series). Thus, potentially, there are important solutions that are missed.

Over the past decade, we have shown optimal control theory to be an efficient and powerful method that can be applied to a wide range of challenging NMR pulse design problems without restricting the space of possible solutions [55–73]. It is capable of designing broadband pulses [66] and selective pulses [74,75] that are simultaneously tolerant to RF inhomogeneity and relaxation effects, which we develop further in the present work.

2. Selective pulse design

Optimal control (including similar, related optimization procedures) was originally introduced into magnetic resonance for the design of band-selective pulses, primarily for imaging [76–82]. It was quickly supplanted by the efficient Shinnar–LeRoux (SLR) algorithm [17–21], which establishes a correspondence between frequency-selective pulse design and digital filter design. There are fast, non-iterative algorithms for the ideal filter and, hence, the ideal pulse. Unfortunately, the algorithm does not accommodate additional criteria, such as tolerance to RF inhomogeneity (included in some of the earliest optimal control-related approaches [76,80]) or relaxation effects. In addition, the most applicable and widely used form of the algorithm derives pulses which produce a specific linear phase dispersion in the spectral response. Pulses producing no phase dispersion, suitable for spectroscopy, are more problematic for the SLR algorithm.

We first provide an overview of well-known issues relevant to selective pulse design, since there is considerably less freedom in the choice of parameters compared to broadband pulses. For

* Corresponding author. Fax: +1 937 775 2222.

E-mail addresses: thomas.skinner@wright.edu (T.E. Skinner), glaser@tum.de (S.J. Glaser).

example, in designing broadband pulses, we have shown [60] there is at best only a marginal relation between the maximum amplitude, $RF_{\max}$, of a pulse and the achievable excitation or inversion bandwidth, as long as the pulse length, T_p , is allowed to increase sufficiently. Alternatively, increasing $RF_{\max}$ for a given T_p can improve performance for a given bandwidth or increase the bandwidth. There can also be innumerable broadband pulses that provide approximately equal performance for a given $RF_{\max}$, T_p , and bandwidth.

Selective pulses are far more constrained, with a well-known relation between the selective bandwidth and T_p , and much tighter limits on the choice of $RF_{\max}$ for a given product of bandwidth and T_p [21]. We provide only a simple overview of the optimal control methodology, emphasizing the modifications necessary for the present work. The basic algorithm for optimizing pulse performance over a range of resonance offsets and RF inhomogeneity is described fully in [57]. Details related to incorporating relaxation [66] and specific dispersion in the phase of the final magnetization [61,67,83], which we refer to as the phase slope, are provided in the associated references.

2.1. Phase slope

Values of the phase slope, R , at each offset [67] characterize the net phase dispersion that accumulates during a pulse of length T_p . The phase slope is defined relative to the net precession of transverse magnetization that would be produced solely by chemical-shift evolution during the same time interval, T_p . A pulse that produces focused magnetization of fixed phase for all spins in the offset range of interest would have constant $R = 0$ (i.e., a self-refocused pulse). Many selective pulses are symmetric, $R = 1/2$ pulses [7,78,21,27,29,30]. The symmetry of the resulting pulse provides an advantage in the development of various algorithms used in selective pulse design, such as SLR, inverse scattering [22,27], polychromatic [29], and stereographic projection [30]. In fact, the standard form of the SLR algorithm [21] can only generate linear phase of this value. Sophisticated algorithms allowing for more general phase in selective pulses are described in the literature [26,84,85], but they are specific to this particular performance factor and cannot include tolerance to variations in other experimentally important parameters.

By contrast, including additional performance criteria, such as a general phase response, is straightforward for optimal control. Initial magnetization $\mathbf{M}(t_0)$ is driven by the RF controls to a final magnetization $\mathbf{F}$ that is defined for each resonance offset in the desired range. To excite transverse magnetization of linear phase slope R , we consider target states for each offset ω in the excitation bandwidth of the form [67]

$$\mathbf{F} = [\cos(\varphi), \sin(\varphi), 0] \quad (1)$$

Choosing $\varphi = R\omega T_p$ gives a linear phase slope, but any function can be considered to define a useful target phase, such as quadratic or higher order.

Selective excitation most simply requires changing the target to $\mathbf{F} = [0, 0, 1]$ for offsets outside the desired bandwidth. In principle, this stopband includes an infinite range of frequencies that must therefore be truncated at some chosen value. We found as a practical matter that choosing the stopband to be ~ 5 times the passband width was sufficient to eliminate excitation at higher frequencies for the pulse parameters used here. This value can easily be adjusted upwards if necessary, or, alternatively, high-frequency components of the resulting pulse determined from Fourier analysis can be deleted after verifying they have no significant effect on the passband excitation.

In addition, since the ideal rectangular offset response cannot be excited by a finite-length pulse, there must be a transition

connecting the excitation frequencies to the nulled frequencies. The selective response profile is typically defined in terms of design parameters for finite impulse response (FIR) filters. An overview of the issues relevant to our optimal control implementation follows.

2.2. Selective pulses as digital filters

For design conditions employing ideal RF in the absence of relaxation, selective pulse performance is completely determined by the desired passband width B , pulse length T_p , transition width W joining the passband and stopband, and residual signal fluctuation or ripple δ_1 and δ_2 about the ideal target amplitude for the passband and stopband, respectively.

The passband frequency ν_p and stopband frequency ν_s are defined where the magnitude of the magnetization response becomes less than the associated fluctuations $1 - \delta_1$ and $|\delta_2|$, as illustrated in Fig. 1 (adopted from Ref. [21]). The frequency where the amplitude drops to one-half is approximately the average of these two frequencies. The full width of the filter is defined as twice this value, giving a bandwidth $B = \nu_s + \nu_p$ and a fractional transition width $W = (\nu_s - \nu_p)/B$.

More specifically (and again emphasizing the design conditions stated at the beginning of the section), selective pulse performance is constrained by relations for optimal FIR filters of the form

$$WT_p B = f(\delta_1, \delta_2), \quad (2)$$

in terms of an empirically derived function $f(\delta_1, \delta_2)$ [86]. For a given value of $f = WT_p B$, smaller (larger) δ_1 gives larger (smaller) δ_2 . Alternatively, for fixed δ_1 or fixed δ_2 , values of f increase as δ_2 or δ_1 , respectively, decrease. Flexibility in selective pulse design is thus purchased at the cost of trade-offs among bandwidth, pulse length, transition width, and ripple amplitudes. Choosing any four of the set determines the fifth.

This relation appears to have been little used in the spectroscopic community. Although the precise form of the function $f(\delta_1, \delta_2)$ holds only for $R = 1/2$ pulses, we have found it to be a useful qualitative indicator for more general R . One important implication is that pulse performance for a given absolute transition width $BW = \nu_s - \nu_p$ can be made independent of the passband width, B . Fixed T_p results in the same performance in terms of residual signal (ripple) for different B as long as the transition width BW is constant. This was observed empirically and noted in [47]. We thus use Eq. [2] to inform our optimal control design.

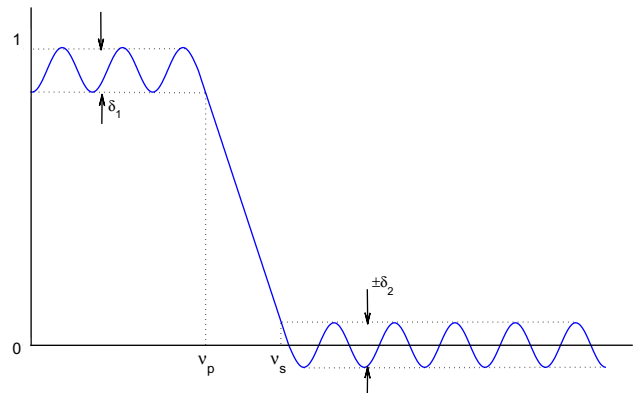


Fig. 1. (Adapted from Ref. [21]) a finite length selective pulse can only approximate the ideal rectangular frequency response. Residual signal or ripple amplitude in the selected frequency spectrum (passband) is denoted by δ_1 , with δ_2 representing the ripple over the frequency range where the signal should be nulled (stopband). The positive frequencies ν_p and ν_s define the passband and stopband, respectively. The plotted response is symmetric about the zero frequency.

Download English Version:

<https://daneshyari.com/en/article/5406165>

Download Persian Version:

<https://daneshyari.com/article/5406165>

[Daneshyari.com](https://daneshyari.com)