

First draft June 2018
Augmented April 2021
Re-augmented March 2025
Re-re-augmented July 2026 (twice)
This draft 8th July 2026

A Perturbative Vade-Mecum

Or: A Technical Note on Strict Perturbation Theory for Nonperturbative Algorithms

Harald W. Griebhammer ^{a1}

^a *Institute for Nuclear Studies, Department of Physics,
The George Washington University, Washington DC 20052, USA*

Abstract

This is a collection of techniques and tricks to perform “strict perturbation theory” around a nonperturbative leading-order solution of various equations, such as integral equations ($T = V + KT$), differential and eigenvalue equations ($HT = ET$ and $HT = T$). Many of the results are not new, but they may appear as a bundled package in an EFT context here for the first time. Much of this information is strewn over the literature, often not in stand-alone articles but as short remarks or appendices in presentations of different foci, with rationales, intermediate steps or comments sometimes at best alluded to.

These notes mark an effort to collect, concisely present and elucidate techniques, similar in spirit to a review. I would therefore highly appreciate if they could be referred to in any publication for which they were deemed useful.

Disclaimer customary for informal drafts: The is is typo-ridden document with huge swaths of non-scientific language, repetitions (sometimes even verbatim), and poor organisation of contents. I have not taken any effort to augment references beyond minimum anecdotal citations, and therefore apologise to all the community that I have overlooked your work.

Suggested Keywords: Effective Field Theory, Perturbation Theory around non-perturbative leading order, parameter fixing, integral equations, differential equations, eigenvalue and eigenvector determination

¹Email: hgrie@gwu.edu

Contents

1	Introduction	2
2	Formalism	3
2.1	Outline	3
2.2	$T = V + VGT$: Operator Equations of the Integral Type With Inhomogeneity	4
2.2.1	Method	4
2.2.2	Normalisation and Convolution	7
2.3	Efficient Fitting Procedures For Any Perturbation	8
2.3.1	Readjusting Parameters from Order to Order: Linear Case	8
2.3.2	Readjusting Parameters from Order to Order: Nonlinear Case	10
2.3.3	Advantage of Order-By-Order Fitting	11
2.3.4	Natural-Sized Coefficients – Or Not?	11
2.4	Further Notes on Mathematical Perturbation Theory	12
2.4.1	Mathematical vs. EFT Expansion	12
2.4.2	Several Expansion Parameters	13
2.4.3	Beyond Scattering	14
2.4.4	Beyond Static Potentials	14
2.4.5	History Lesson	14
2.5	$ET = KT$: Eigenvalue Equations with Varying Eigenvalue	15
2.5.1	Solving Without Resort to a Complete Ortho-Normal System	16
2.5.2	Solving Via A Complete Ortho-Normal System of Eigenfunctions to LO	19
2.6	$T = KT$: Eigenvalue Equations with Fixed Eigenvalue	21
2.7	Distorted Wave Born Approximation	23
2.8	More Variations	24
3	From Solutions to Observables	24
3.1	Unitarity	24
3.2	Extracting Phase Shifts	25
3.3	Finding and Moving Poles in Amplitudes	26
3.3.1	Method	26
3.3.2	Example: NLO Effective Range Theory	27
3.3.3	Example: Zeeman and Paschen-Back Effect in Hydrogen	31
4	Summary and Conclusions	31
A	Future Appendices	32
	Bibliography	32

1 Introduction

In χ EFT, EFT(π) and many other descriptions, the dominant part of a physical process involves often an infinite number of contributions by iterating a few interaction over and over again, e.g. One Pion Exchange (OPE) and a few contact interactions. Interactions at next-to-leading order (NLO) and higher (NⁿLO) are corrections to the LO result.

“Strict” perturbation about a leading-order (LO) result which is non-perturbative has many advantages, including the absence of spurious bound states (beyond those found already at LO), simpler fitting routines with less and more pronounced minima (since coefficients appear as simple polynomials, not inside nonlinear/nonperturbative dependencies), simpler implementation of few-body interactions, applicability of a wider range of “cutoffs” to regularise the interactions and hence better control in answering questions about whether a theory is renormalised at a given order, etc. Foremost, however, “higher-order” corrections must be small relative to LO, i.e. resumming them cannot significantly change the LO result. If the correction is as sizeable or even larger than LO, then LO is not LO.

Another advantage is that LO is often very simple. In EFT(π), for example, LO interactions involve 2- and 3-body interactions which do not mix partial waves: There is only one result for each value of spin and orbital angular momentum, independent of the total angular momentum to which they are coupled. That means the LO solution is numerically very simple and requires very little storage space, and book-keeping of partial-wave mixing is only included as perturbation.

The techniques described here were not invented by me; but I have added a few extensions and flourishes in places. I am bundling information from several sources, including deep, dark pockets of my brain where no proper list of references is maintained. I am most likely to overlook original work which should be cited, as well as not even making a *bona-fide* effort of tracking down the first appearance in the literature.

The benefit of these notes may thus hopefully lie in the fact that they are provided with the goal of being written pedagogically, in a unified notation, and with several methods congregated in one place – similar to a review. I would therefore highly appreciate if readers could refer to it in any publication in which it was deemed useful.

Math aficionados will notice that many equations can be compressed quite a bit by using Dirac’s bra-ket notation. To employ one of the all-time favourite quotes of the lazy person: “This is left as an exercise to the reader.”

The notes are organised as follows. Section 2 is focused on the application of Mathematical Perturbation Theory to the type of problems usually found in Physics. After a brief overview of the method (sect. 2.1), it addresses first operator equations of the integral type With inhomogeneity, *i.e.* Lippmann-Schwinger-like problems $T = V + VGT$ (sect. 2.2, followed by two sub-sections with remarks which are generally applicable in perturbative settings, like parameter fitting (sect. 2.3). It then turns to Eigenvalue Equations with fixed eigenvalue, $T = KT$ (sect. 2.6), those with varying eigenvalue like Schrödinger-like equations (sect. 2.5), and possibly other equations.

Section 3 discusses how to extract several observables from formal solutions, starting

with tests of unitarity (sect. 3.1). Examples include phase shifts (sect. 3.2) and evaluating shifting bound-state energies in perturbation (sect. 3.3), but most of this is kept at a simple level, pointing to the literature for formulae involving more complicated cases like composite systems or coupled channels.

Customary conclusions and outlook follow in sect. 4.

2 Formalism

2.1 Outline

I am quite certain that all I am writing here has already been discussed in the extensive literature on Mathematica Perturbation Theory, see e.g. [1–3] for excellent, brief and cheap accounts. Much of this and some more is also reflected in lectures on Mathematical Methods for Theoretical Physics at George Washington University [4].

Mathematical Perturbation Theory deals with problems of the form

$$F[T; K; C] = 0 \quad , \quad (2.1)$$

where T is some set of unknowns, K some set of operators for which (hopefully) a perturbative expansion exists, and C some set of parameters which also (hopefully) permit an expansion in a small parameter. Sometimes, we will use $VG = K$ instead of K . The theory now assumes at the start that a common, small expansion parameter ϵ exists in which *all* variables can be expanded, represented symbolically as [In this notation, the subscript itself determines the order.]:

$$\begin{aligned} T &= \epsilon^0 T_0 + \epsilon^1 T_1 + \epsilon^2 T_2 + \dots \\ K &= \epsilon^0 K_0 + \epsilon^1 K_1 + \epsilon^2 K_2 + \dots \\ C &= \epsilon^0 C_0 + \epsilon^1 C_1 + \epsilon^2 C_2 + \dots \end{aligned} \quad (2.2)$$

$F[T; K = VG; C] = 0$ is then expanded in that series, and since the expansion parameter is assumed to be small but arbitrary, each order in ϵ must be zero by itself. One now solves the equations by induction from the lowest order to the highest, with each lower-order result feeding into higher-order equations. I will clarify this abstract description a bit more concretely in the first example, sect. 3.3.1.

Comment on the Mathematical Expansion Parameter Very often, ϵ is only used as “tracking device” and actually sets $\epsilon \rightarrow 1$ in the end while the terms $\epsilon^i V_i \rightarrow V_i$ and $\epsilon^i G_i \rightarrow G_i$ by themselves are assumed to be a successively smaller and smaller. That is for example the case when one considers perturbations of matrices.

Operators All operators are assumed to be self-adjoint, and all solutions to be normalisable in their Hilbert space, at least in the distributional sense. That implies that every operator equation has a Complete Ortho-Normal System (CONS) of solutions which can be used as basis of the Hilbert space on which it acts.

Checking Assumptions Mathematical Perturbation Theory comes with a consistency check: The expansion is performed under the *assumption* that “higher-order terms” provide ever-decreasing contributions. After the calculation has been performed, it is absolutely crucial to always check that this is indeed the case. If higher-order corrections are not small, then the assumption has failed. This can be quantified and turned into a semi-quantitative criterion of probabilities of failure in a Bayesian analysis of order-by-order convergence; *cf.* the work of the BUQEYE group around D. R. Phillips and R. Furnstahl and of the BAND collaboration.

The Simple Case Only I only report on *Non-Degenerate/Regular* Perturbation Theory, *i.e.* I assume throughout that each order in perturbation has one and only one solution. If one is faced with a singular/degenerate problem (*e.g.* more-than-one-dimensional eigenspaces to the same eigenvalue), one can fortunately often still use non-degenerate/regular perturbation theory. Rephrasing Griffith/Schroeter’s “Moral” [5, sec. 7.2], this works when all of the following criteria are met:

- (1) All degenerate unperturbed solutions can be labelled by what Physicists call additional quantum numbers to a set of operators $\{\mathcal{O}_\alpha\}$, $\alpha = 1, \dots$
- (2) This set is mutually commuting and commutes with the perturbation operators up to the desired order.
- (3) Each state is uniquely identified by a label which contains the quantum numbers to all operators $\mathcal{O}_\alpha$.

If all this fails, you have my sympathies and must resort to the “method of undetermined gauges” of **Singular/Degenerate Perturbation Theory** described *e.g.* in Murdock [2, sect. 1.7]; see also *e.g.* [1, 4].

2.2 $T = V + VGT$: Operator Equations of the Integral Type With Inhomogeneity

2.2.1 Method

Let us start with the simplest case, namely an operator equation

$$T = V + VGT \tag{2.3}$$

where V and G is given and T is to be solved for. This type appears as integral equations (scattering, Lippmann-Schwinger, Faddeev, Skorniakov-Martirosian and other few- and many-body equations), but many other equations can be written this way as well. For the sake of concreteness, I use the following nomenclature, but other names are used in other descriptions/equations which reduced to operator equations of the same type:

- V is the *potential*;

- G is the *propagator* (e.g. for A particles, often nonrelativistic);
- T is the *T -matrix/scattering amplitude*;
- the first term on the right (the one without T) is the *driving term* or *inhomogeneity*;
- the second term is the *homogeneous term* and contains the *kernel* $K = VG$.

Let us now assume that an expansion in a small parameter $\epsilon \ll 1$ exists such that

$$V = V_0 + \epsilon^1 V_1 + \epsilon^2 V_2 + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n V_n \quad (2.4)$$

$$G = G_0 + \epsilon^1 G_1 + \epsilon^2 G_2 + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n G_n \quad (2.5)$$

An expansion for the potential is exactly what EFT provides (but see note below for the difference between the mathematical expansion parameter and the EFT expansion parameter Q). It also provides an expansion for G , e.g. by relativistic corrections to non-relativistic propagators via suppression in powers of p/M (typ. momentum scale over typ. particle mass). Often, that implies that V_n is nonzero right away, while G_n carries several corrections $G_1, G_2 \equiv 0$ which are zero. For simplicity and because it covers the most interesting case already, I will therefore assume that there is actually no expansion for the propagator, i.e. $G \equiv G_0$; see additional note and generalisation below.

Mathematical Perturbation Theory asserts that there exists a likewise expansion for

$$T = T_0 + \epsilon^1 T_1 + \epsilon^2 T_2 + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n T_n \quad (2.6)$$

when the solution T is not degenerate at leading order; see remark on Singular/Degenerate Perturbation Theory in sect. 2.1.

Following Mathematical Perturbation Theory, one now inserts eqs. (2.4) to (2.6) into eq. (2.3) and equates like powers of ϵ . The idea is simple: the equation is to hold for any, arbitrary $\epsilon \ll 1$ – so while there could be particular values of ϵ which allow for solutions in which different orders of ϵ conspire, such fine-tuning is not the generic case we are after. One then finds

$$\begin{aligned} \mathcal{O}(\epsilon^0) : T_0 &= V_0 + V_0 G T_0 = V_0 (\mathbb{1} + G T_0) \\ \mathcal{O}(\epsilon^1) : T_1 &= V_1 (\mathbb{1} + G T_0) + V_0 G T_1 \\ \mathcal{O}(\epsilon^2) : T_2 &= V_2 (\mathbb{1} + G T_0) + V_1 G T_1 + V_0 G T_2 \\ \mathcal{O}(\epsilon^n) : T_n &= V_n (\mathbb{1} + G T_0) + \sum_{m=1}^{n-1} V_{n-m} G T_m + V_0 G T_n \end{aligned} \quad (2.7)$$

[Notation: The subscripts of symbols which are multiplied with each other sum to the order n considered. That also implies that each term in a sum must have the same order. Example: $V_1 G T_1$ has $n = 1(V) + 1(T) = 2$ and is hence N²LO, as is $V_0 G T_2$ to which it is added.]

At leading order, we recover the iterative solution T_0 . It is multiplied with the LO potential V_0 and fed into the NLO integral equation as part of the driving term $V_1(\mathbb{1} + G)T_0$, and one solves the integral equation for T_1 . This continues at N²LO: the solutions $T_{0,1}$ of the lower-order equations are fed into the driving terms of the N²LO integral equation, etc. Because of the geometric form eqs. (??) to (2.7) take, with a slim top and a bulky bottom, I call this a *pyramid* (not a pyramid scheme).

The kernel of the integral equation (see last part in each line) is at each order given by V_0G , so that the homogeneous integral equation is always the same. Only the inhomogeneity gets more complicated with each order, but is always constructed from the potentials and solutions $V_{i < n}, T_{i < n}$ of the preceding orders. Therefore, one can recycle the computational method to solve the LO equation at every higher order by simply changing the inhomogeneity, but not the kernel, with $\mathbb{1} - V_0G$ the only operator ever to be inverted:

$$\begin{aligned}
\mathcal{O}(\epsilon^0) : T_0 &= [\mathbb{1} - V_0G]^{-1} V_0 \\
\mathcal{O}(\epsilon^1) : T_1 &= [\mathbb{1} - V_0G]^{-1} V_1(\mathbb{1} + GT_0) \\
\mathcal{O}(\epsilon^2) : T_2 &= [\mathbb{1} - V_0G]^{-1} [V_2(\mathbb{1} + GT_0) + V_1GT_1] \\
\mathcal{O}(\epsilon^n) : T_n &= [\mathbb{1} - V_0G]^{-1} \left[V_n(\mathbb{1} + GT_0) + \sum_{m=1}^{n-1} V_{n-m}GT_m \right]
\end{aligned} \tag{2.8}$$

[It can be computationally advantageous to not invert but solve the linear system of equations (??-2.7); that depends on context.] As a result, this algorithm by induction produces in one sweep all orders up to and including the desired one as a natural byproduct. That also means Bayesian error estimates (“max-criterion”) can immediately be applied (possibly after some reordering; see note below on the difference between the mathematical expansion parameter ϵ and the EFT expansion parameter Q).

The inverse $[\mathbb{1} - V_0G]^{-1}$ can be computed once and stored. Since LO interactions are usually also quite simple, the inversion can be performed with very high accuracy. For example in EFT(π), LO interactions are S-waves only between 2 and 3 particles, respectively, and are furthermore spin-independent. The range of quantum numbers is thus very small. The CPU complexity saved by book-keeping fewer quantum numbers can be invested in finer momentum grids in both T_0 and $[\mathbb{1} - V_0G]^{-1}$.

This can be important because the induction from one order to the next implies that small numerical uncertainties in low-order solutions like T_0 and in $[\mathbb{1} - V_0G]^{-1}$ can become exponentially large when inserted in high-order insertions if each order’s numerical errors are uncorrelated with the next. For example, an N³LO amplitude T_3 comes from 3 insertions of T_0 and $[\mathbb{1} - V_0G]^{-1}$, on top of the LO solution. One can thus assume the error will be increased at least with a 3rd power. This is the reason why it is wise to invest the reduction of numerical complexity from a “simple” LO interaction into a finer momentum grid.

If G needs to be expanded as well, then this reshuffles terms a bit:

$$\begin{aligned} \mathcal{O}(\epsilon^n) : T_n &= V_n(\mathbb{1} + G_0 T_0) + \sum_{m=1}^{n-1} \left[\sum_{m'=0}^m V_{m'} G_{n-m-m'} \right] T_m + V_0 G_0 T_n \\ \implies T_n &= [\mathbb{1} - V_0 G]^{-1} \left[V_n(\mathbb{1} + G_0 T_0) + \sum_{m=1}^{n-1} \left[\sum_{m'=0}^m V_{m'} G_{n-m-m'} \right] T_m \right] \end{aligned} \quad (2.9)$$

2.2.2 Normalisation and Convolution

The eigenvector $T = T_0 + T_1 + \dots$ is not normalised. If one chooses to do that, then the normalisation must be truncated *consistent with the order-by-order expansion*; *e.g.* at NLO:

$$|T|^2|_{\text{NLO}} = T_0^\dagger T_0 + T_0^\dagger T_1 + T_1^\dagger T_0 = T_0^\dagger T_0 + 2\text{Re}[T_0^\dagger T_1] + \mathcal{O}(\epsilon^2) \stackrel{!}{=} 1 \quad (2.10)$$

or in general

$$|T|^2|_{\text{N}^n\text{LO}} = \sum_{0 \leq m+l \leq n} T_m^\dagger T_l \stackrel{!}{=} 1 \quad (2.11)$$

Including more terms (*e.g.* $T_n^\dagger T_n$) does not increase the accuracy of the expansion since other terms from interactions V_{n+1} *etc.* are absent. It is also simply inconsistent with the tenets of Mathematical Perturbation Theory.

When T is the T -matrix of a scattering problem, the normalisation is slightly different. I will use a very simple relation between T - and S -matrix: $S = \mathbb{1} + 2ik k T$, skipping all factors. Since the S -matrix is unitary, one performs again the expansion consistently up to N^nLO . Below the first inelasticity, this gives:

$$1 \stackrel{!}{=} \left| \mathbb{1} + 2ik \sum_{m=0}^n T_m \right|^2 \Big|_{\text{N}^n\text{LO}} \implies \sum_{m=0}^n \text{Im}[T_m] \stackrel{!}{=} k \sum_{0 \leq m+l \leq n} T_m^\dagger T_l \quad (2.12)$$

The last equality is simply the statement that the Optical Theorem must hold order-by-order. Above the first inelasticity, the conditions are replaced by $\stackrel{!}{\leq}$ if not all open channels are described.

The degree to which these equations are numerically satisfied is a criterion of numerical accuracy and of the convergence of the perturbative series. Remember that the expansion is performed under the *assumption* that “higher-order terms” provide ever-decreasing contributions. One must always check that this is indeed the case.

Likewise, observables produced by convoluting T with operators $\mathcal{O}$ should be expanded consistently. I will at some point write a Section 3 about that, but for now I refer to a recent article whose Appendix A details the order-by-order expansion of phase shifts also in coupled channels and binding energies from the T -matrix [6] and to an application in Chiral EFT with complementing information [7, 8]; see also sect. 3.3.

Converting the results of Mathematical Perturbation Theory to observables, including further comments on unitarity, will be addressed in sect. 3.

One could now turn to addressing perturbations for the bound-state solution to this type in problems, $T = KT$. However, we consider first eigenvalue equations of the Schrödinger type. Combining these with the nonlinear fit procedure described in sect. 2.3.2 will provide a general solution strategy for $T = KT$ in sect. 2.6. Even before that, I will add sub-sections useful for Physics practitioners.

2.3 Efficient Fitting Procedures For Any Perturbation

This note applies actually to all techniques discussed in this article, not only to operator equations of the integral type. It should be re-written to be shorter.

In EFT, short-distance parameters appear which need to be determined from some input, e.g. data or lattice QCD. That appears to produce a problem. For argument's sake, perform a LO calculation with a fit parameter C and only one new structure D at NLO which requires fitting. The problem is that our LO solution $T_0(C_{\text{LO}})$ depends now on the LO value of C , while the NLO fits C_{NLO}, D enter when we want to determine $T_1(C_{\text{NLO}}, D)$. To find $T_1(C_{\text{NLO}}, D)$, we cannot use the LO solution $T_0(C_{\text{LO}})$ because it depends on a different value $C_{\text{LO}} \neq C_{\text{NLO}}$. We appear to need to produce a *new* LO result $T_0(C_{\text{NLO}})$ with an adjusted NLO C_{NLO} which we can feed into the NLO equation to get the NLO solution. Without really being able to recycle the LO calculation which used the LO value for C , the method above loses much of its appeal. And fitting even at NLO becomes cumbersome since we need to run the pyramid from top to bottom for every single fit iteration. Imagine that for an N^4LO calculation.

The solution is actually again given in the literature of Mathematical Perturbation Theory: Along with the operators of the equations, the coefficients also *must* be expanded in a power series in ϵ . This is often discussed in connection with the anharmonic oscillator or φ^4 -corrections which turn the harmonic oscillator into a pendulum [4]. To my knowledge, this trick has first been used in nuclear EFT by Mehen and Stewart [14] in the context of the KSW approach, and quickly became an instrumental technique in making higher-order results amenable and thus easier, especially in $\text{EFT}(\not{\pi})$.

So we write

$$C = C_0 + \epsilon^1 C_1 + \epsilon^2 C_2 + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n C_n \quad (2.13)$$

$$D = \left[\epsilon^0 D_0 + \epsilon^1 D_1 + \mathcal{O}(\epsilon^2) \right] = \sum_{n=0} \epsilon^n D_n \quad (2.14)$$

for (in principle) all coupling constants (but in practise, only for those we fit – the others only need readjusting if UV divergences need to be absorbed; see below). It may seem somewhat strange that the expansion of both coefficients C and D starts at ϵ^0 , but that is just a notational choice which becomes clear in a moment.

2.3.1 Readjusting Parameters from Order to Order: Linear Case

Let us first assume that our potentials have the form $V_0 = C\hat{V}_0$ at LO, $V_1 = D\hat{V}_1$ at NLO, $V_2 = E\hat{V}_2$ at N^2LO etc., where the “hat” indicates the operator structure only (no parameter

to be adjusted/determined is hidden inside it). We now just need to re-order the potentials by collecting all terms which actually contribute at a given order in ϵ and replace the terms V_i in eqs. (??) to (2.7) by the new power-counted potentials:

$$\begin{aligned}
\mathcal{O}(\epsilon^0) : V_0 &\rightarrow C_0 \hat{V}_0 \\
\mathcal{O}(\epsilon^1) : V_1 &\rightarrow C_1 \hat{V}_0 + D_0 \hat{V}_1 \\
\mathcal{O}(\epsilon^2) : V_2 &\rightarrow C_2 \hat{V}_0 + D_1 \hat{V}_1 + E_0 \hat{V}_2 \\
\mathcal{O}(\epsilon^3) : V_n &\rightarrow \sum_{m=0}^n \mathcal{C}_{nm} \hat{V}_{n-m} .
\end{aligned} \tag{2.15}$$

The last line denotes the generic coefficient of the N^n LO potential as $V_n = \mathcal{C}_n \hat{V}_n$ and expands $\mathcal{C}_n = \sum_{m=0}^n \mathcal{C}_{nm}$ as for C, D, E . [The choice of notation for the subscript of the coefficients guarantees now that we can just sum subscripts in products to determine an overall order, and to quickly check that sums of contributions have the same order.] Indicating the dependence of the solutions T_i on the coefficient C_i, D_i , the pyramid of equations becomes

$$\begin{aligned}
\mathcal{O}(\epsilon^0) : T_0(C_0) &= C_0 \hat{V}_0 + C_0 \hat{V}_0 G T_0(C_0) \\
\mathcal{O}(\epsilon^1) : T_1(C_0, C_1; D_0) &= [C_1 \hat{V}_0 + D_0 \hat{V}_1][\mathbb{1} + G T_0(C_0)] + C_0 \hat{V}_0 G T_1(C_0, C_1; D_0) \\
&\text{etc.}
\end{aligned} \tag{2.16}$$

[I am not recording the general case since the equations become lengthy, without a lot of information content. “This is left as an exercise to the reader.”]

Now, one *can* recycle the LO calculation and its fit (which determined C_0) in the NLO calculation. Along with T_0 , the value C_0 is fed into the next-order equation as known, and C_1, D_0 are fit parameters at NLO. The price we appear to pay is that we now have to fit 2 new parameters at NLO: C_1 and D_0 , while we only had 1 parameter D before. Up to N^2 LO, 5 parameters C_0, C_1, C_2, D_0, D_1 appear instead of 2 (C, D). Do we need one more datum at NLO, plus 2 more at N^2 LO?

Of course not. Recall that we already used C_0 at LO to reproduce a datum, say a binding energy. That the value of that observable shall not change at NLO provides an additional constraint and brings us back to needing only one new datum to adjust D_0 at NLO. Usually, there is no bijection between observable and parameter (except if you choose to parametrise your potential in a clever way, see e.g. [15] for an example to fit the momentum-dependent 3-nucleon interaction H_2 at N^2 LO in EFT(π)). One then ends up with a linear system of two coupled equations for two unknowns at NLO: One constraint that the LO datum shall still be reproduced at NLO, and one new datum to be fit at NLO. At N^2 LO, we get 2 constraints: both data we already fitted should not move.

For more fit parameters and at higher orders, the same procedure applies accordingly. At each order, one obtains fit equations which are a combination of “lower orders shall not move” and “reproduce this datum”.

It is worth stressing again that this is at each order a linear system of equations to fit, that means there is one and only one unique solution which on top of that is computed quite

straightforwardly by matrix inversion even in problems with very many parameters. This stands in contrast to approaches in which the problem is solved non-perturbatively,

$$\begin{aligned} T &= (C\hat{V}_0 + D\hat{V}_1 + \dots) + (C\hat{V}_0 + D\hat{V}_1)GT \\ \implies T &= [\mathbb{1} - (C\hat{V}_0 + D\hat{V}_1)]^{-1}C\hat{V}_0 + D\hat{V}_1 . \end{aligned} \quad (2.17)$$

At each order, the inversion now depends on each parameter, and parameter determinations become highly nonlinear problems in multi-dimensional spaces. It becomes notoriously difficult to find solutions, and even if one does, showing that a given solution optimises the penalty function of the fit (“global minimum”) is challenging. So, this is another advantage of using perturbation theory for higher-order calculations.

2.3.2 Readjusting Parameters from Order to Order: Nonlinear Case

Now, let us discuss the case the the potentials C_n depend highly non-linearly on some parameter which needs to be fitted. Concretely, one could consider fitting the pion mass in the one-pion and two-pion exchange potentials from data (which could be useful to check underlying Physics – Nijmegen has done this for the 2PE). The pion mass enters highly non-linearly in amplitudes (and in observables).

Still, the contribution involving C_1 becomes NLO, so that we can rewrite the potential by expanding the potential in a series around C_0 :

$$\begin{aligned} V_0(C) + \epsilon V_1(D) &= V_0(C_0) + \epsilon^1 C_1 V'_0(C_0) + \epsilon^2 \left[C_2 V'_0(C_0) + \frac{1}{2} C_2^2 V''_0(C_0) + \mathcal{O}(\epsilon^3) \right] \\ &\quad + \epsilon^1 V_1(D_0) + \epsilon^2 D_1 V'_1(D_0) + \epsilon^2 V_2(\dots) + \mathcal{O}(\epsilon^3) \\ &= V_0(C_0) + \epsilon^1 \left[C_1 V'_0(C_0) + V_1(D_0) \right] \\ &\quad + \epsilon^2 \left[C_2 V'_0(C_0) + \frac{1}{2} C_1^2 V''_0(C_0) + D_1 V'_1(D_0) + V_2(\dots) \right] + \mathcal{O}(\epsilon^3) \end{aligned} \quad (2.18)$$

where derivatives are with respect to the parameter(s) of the potential. [Notice that “ $V_1(D_0)$ ” denotes a term of the potential which only enters at ϵ^1 and is absent at order ϵ^0 . This notation allows me to still determine the order at which a contribution enters by adding the subscripts. A previous version of these notes had a different notation which was more convoluted.]

You see that the higher-order corrections to a coupling linearise, even if the parameter enters non-linearly at first order. Counter-terms in EFT are usually linear in the coupling, i.e. $C_1 V'_0(C_0) = V_0(C_1)$, $C_2 V'_0(C_0) = V_0(C_2)$, $C_n V'_0(C_0) = V_0(C_n)$ and $V''_n \equiv 0$.

We now replace the terms V_i in eqs. (??) to (2.7) by the new power-counted potentials

$$\begin{aligned} V_0 &\rightarrow V_0(C_0) \\ V_1 &\rightarrow V_1(D_0) + C_1 V'_0(C_0) \\ V_2 &\rightarrow C_2 V'_0(C_0) + \frac{1}{2} C_2^2 V''_0(C_0) + D_1 V'_1(D_0) + V_2(\dots) \end{aligned} \quad (2.19)$$

Indicating the dependence of the solutions T_i on the coefficient C_i , the pyramid of equations becomes then

$$\begin{aligned}\mathcal{O}(\epsilon^0) : T_0(C_0) &= V_0(C_0) + V_0(C_0)GT_0(C_0) \\ \mathcal{O}(\epsilon^1) : T_1(C_0, C_1; D_0) &= [C_1V'_0(C_0) + V_1(D_0)][\mathbb{1} + GT_0(C_0)] + V_0(C_0)GT_1(C_0, C_1; D_0) \\ \text{etc.}\end{aligned}\tag{2.20}$$

[I am not recording the general case since the equations become lengthy, without a lot of information content. “This is left as an exercise to the reader.”]

As in the linear case, we now *can* recycle the LO calculation and its fit (which determined C_0) in the NLO calculation; *etc.*. The procedure described for the linear case carries through. These equations are nonlinear in C_0 , *i.e.* at LO, but they become polynomial at higher orders: NLO is linear in the first correction C_1 of the LO coefficient. At N²LO, the system is quadratic in C_1 but linear in C_2 and D_0 ; *etc.*

For more fit parameters and at higher orders, the same procedure applies accordingly. At each order, one obtains fit equations which are a combination of “lower orders shall not move” and “reproduce this datum”. From NLO on, these equations involve the fit parameters as polynomials. For a parameter first entering at N^{*m*}LO, the polynomial at N^{*n*}LO ($n \geq m$) is at most of degree $(1 + n - m)$. We can therefore solve these quite quickly.

2.3.3 Advantage of Order-By-Order Fitting

Strict perturbation avoids highly nonlinear fitting procedures, and now we even have nice way to recycle N^{*n*}LO fits in a N^{*(n+1)*}LO calculation, so that running the pyramid just once provides all orders up to the required one in one go, with fits to the requested data at each order. We can recycle all lower-order results when we move to the next order; we do not need to re-run the fitting gauntlet, only at the new order we add. We can actually constrain the fit algorithm to reflect that the higher-order correction $C_1, \dots$ to C_0 should be parametrically small and search for new minima only in proximity, *i.e.* perturbatively close, to the lower-order ones. But see the following caveat:

2.3.4 Natural-Sized Coefficients – Or Not?

That, however, leads to a word of warning. The procedure assumes that the fit parameters obey a perturbative expansion, with $C_0 \gg C_1 \gg C_2 \gg \dots$, *etc.* This certainly holds for *renormalised* parameters. If some C_i is strongly cutoff dependent (“diverges in the UV”, like most bare couplings), then this will not hold, even though the NLO amplitude T_1 is perturbatively close since the divergences between this and other terms cancel. This situation is familiar from QED, where bare couplings may diverge, $g^2/(d-4)$ in dimensional regularisation with $d \rightarrow 4$. In everyday QFT, we often still treat such a coupling as “parametrically small” (basically say $g \rightarrow 0$ quicker than $d \rightarrow 4$ because calculations are analytic and results can be analytically continued to the real world where $g \nrightarrow 0$). But we do not have that luxury when cutoffs are employed. We might get away with it, but there is no

guarantee. We might have to pay a price by choosing cutoffs around the breakdown scale $\Lambda \sim \bar{\Lambda}_{\text{EFT}}$ so that the “natural” ordering $C_0 \gg C_1 \gg C_2 \gg \dots$ holds. Test and see.

None of this is actually exciting news for those familiar with renormalisation. At higher orders, coefficients are re-normalised, i.e. their strengths get adjusted by parametrically small amounts. Renormalised couplings have an expansion in small parameters, bare ones do not.

2.4 Further Notes on Mathematical Perturbation Theory

After this first example, we also understand some comments which also apply beyond the integral-equation case.

2.4.1 Mathematical vs. EFT Expansion

The mathematical expansion parameter ϵ is *not* the EFT expansion parameter Q (typical low-momentum scales like the pion mass and the scattering momentum, in units of the breakdown scale $\bar{\Lambda}_{\text{EFT}}$).

First, remember that in order to make the summation of an infinite number of interactions mandatory at LO, one can easily derive a consistency equation for the power counting [13]: When some interactions must be treated non-perturbatively at leading order because of shallow real or virtual bound states with scales $p_{\text{typ}} \ll \bar{\Lambda}_{\text{EFT}}$ in the EFT’s range of validity, all terms in the LO equation (??), including the potential, must be of the same order (when all nucleons are close to their non-relativistic mass-shell). If not, one term could be treated as perturbation of the others and there would be no shallow bound-state: no infinite number of diagrams is summed, and so the system can decompose into non-interacting pieces after a while. Thus, $T_0 \sim V_0$, with an exponent m which is determined by the fact that the homogeneous term also must obey $V_0 G T_0 \sim Q^m$. So m also depends on G .

To be concrete, let us consider NN scattering at nonrelativistic energies, i.e. eq. (??) is the LO Lippmann-Schwinger equation. Picking the nucleon-pole in the energy-integration, $q_0 \sim \frac{q^2}{2M}$ and using Q^3 for the momentum integration as well as $G \sim Q^{-2}$ for the non-relativistic NN propagator leads to the consistency condition $V_0 G T_0 \sim Q^m Q^3 Q^{-2} Q^m \sim Q^{2m+1} \stackrel{!}{\sim} Q^m$, i.e. $T_0 \sim V_0 \sim Q^{-1}$. In particular, the one-pion exchange appears to scale as $(\vec{\sigma}_1 \cdot \vec{q})(\vec{\sigma}_2 \cdot \vec{q})/(\vec{q}^2 + m_\pi^2) \sim Q^0$ when one counts only explicit low-momentum scales, but must be of order Q^{-1} if its iteration is to be mandated, like in χEFT .

In a straightforward extension, amplitude and interaction between n nucleons scale as

$$T_{nN} \sim V_{nN} \sim Q^{1-n} \quad (2.21)$$

if it is nonperturbative at LO.

After that little digression, why is ϵ now different from Q ? As you see, Q is usually negative, while eqs. (2.4) to (2.6) start with a zeroth power for ϵ . But what is more, consider a contribution which is of very high order in the EFT expansion. Correlated Three-Pion Exchange (3PE) in NN is of higher order than the N²LO Two-Pion Exchange (2PE) or the leading short-range corrections. So, while we have to insert 2PE twice or even more

often at a high order ($\epsilon^{\geq 2}$), we will need to insert 3PE only once (ϵ^1) to get the same overall accuracy. While it is possible to reformulate the series such that ϵ merges into Q , I prefer to keep the mathematical and EFT expansion parameters different.

2.4.2 Several Expansion Parameters

If more than one low-momentum scale is present, an expansion in multiple small, dimensionless parameters Q_l may be needed. This is usually a pain in the neck. Instead, one often resorts to exploiting numerical coincidences. Much of the following text is slightly modified from ref. [9] which talks about Compton scattering. The multi-scale method by itself is not new.

Take χ EFT with explicit $\Delta(1232)$ degrees of freedom, where three typical low-energy scales exist: the pion mass m_π as the typical chiral scale; the Delta-nucleon mass splitting $\Delta_M \approx 300$ MeV; and a low-energy momentum scale k at which the process takes place. Each provides a small, dimensionless expansion parameter when measured in units of a natural “high” scale $\Lambda_\chi \gg \Delta_M, m_\pi, k$ at which the theory is expected to break down because new degrees of freedom enter. While a three-parameter expansion is possible, it is more economical to follow Pascualtsa and Phillips [10] and take advantage of a convenient numerical coincidence by identifying

$$\delta \equiv \frac{\Delta_M}{\Lambda_\chi} \approx \left(\frac{m_\pi}{\Lambda_\chi} \right)^{1/2} \approx 0.4 \ll 1 . \quad (2.22)$$

For simplicity, we count $M_N \sim \Lambda_\chi$ and employ one common breakdown scale $\Lambda_\chi \approx 650$ MeV, consistent with the masses of the ω and ρ as the next-lightest exchange mesons.

A note on the relevant Δ scales. If a photon or pion hits the nucleon, then all of its energy is absorbed and the Δ resonance is excited with a peak at $\Delta_M \approx 300$ MeV, providing the scale above. The scales are different for processes like NN scattering where the beam particle cannot be annihilated. In that case, the kinetic energy of the relative motion between the nucleons must be equated, i.e. $\Delta_M \approx p^2/(2M_N)$, so that $p \sim 570$ MeV appears to be a better scale. However, that is not quite the whole story. In the two-pion loop where both pions are on their mass-shell, the intermediate $2N$ state sees the Delta resonance at $\Delta_M \approx 300$ MeV, again because the pions’ energies and momenta are on an equal footing.

Notice also that the $\Delta(1232)$ has a large width $\Gamma \approx 110$ MeV, implying that the effects of the resonance are seen well below an energy $\approx \Delta_M$. The Δ width is not of Breit-Wigner form, but energy-dependent and nonzero only for $E \geq m_\pi$.

In such a multi-scale world, one can then identify two regimes of momenta k where the counting simplifies further [11, 12]. In *regime I*, $k \lesssim m_\pi$ counts like a chiral scale, $k \sim m_\pi \sim \delta^2 \Lambda_\chi \ll \Delta_M$, and pion-cloud physics dominates.

In contradistinction, in *regime II*, $k \sim \Delta_M \sim \delta^1 \Lambda_\chi \gg m_\pi$, and the Delta resonance may be excited so that it strongly dominates the relevant channels and dwarfs contributions from the pion cloud through its large width and strong $\gamma N \Delta$ coupling.

Because of the increasing momentum and the reordering of contributions expected around the Delta resonance, the level of theoretical uncertainty is actually different in these

two regimes. In Compton scattering, $k = \omega$ is set by the photon energy. An amplitude which contains in regime I all contributions at $\mathcal{O}(e^2\delta^4)$ (N⁴LO, accuracy $\delta^5 \approx 2\%$) is accurate in regime II only at $\mathcal{O}(e^2\delta^0)$ (NLO, accuracy $\delta^2 \approx 20\%$). For even higher energies, the expansion parameter approaches one, and the EFT series does not converge.

2.4.3 Beyond Scattering

The formalism is often applied to NN or few-nucleon scattering only, but it can easily be extended to any nuclear reaction, including electron scattering $eX \rightarrow eX$ (form factors), Compton scattering $\gamma X \rightarrow \gamma X$, recombination and breakup. The idea is simple: Insert your interaction together with the NN (and few- N) potential(s) into the equations above and keep the terms which produce final and initial states. I will provide a concrete example when I have time.

2.4.4 Beyond Static Potentials

If V is indeed a potential, then it usually does not include retardation effects or radiative corrections, e.g. of on-shell pion-propagation while nucleons propagate. Including these makes the potential V more complex and may well need additional work.

2.4.5 History Lesson

For decades, EFT practitioners like me refused to go beyond NLO in strict perturbation theory, even in EFT(π). The reason was simple. When one write perturbation theory as Feynman diagrams (which is natural for EFT people who often have an education in Particle Physics), then one contribution to an N²LO calculation mandates two insertions of the NLO contributions, Big-Mac'd between the leading-order amplitude: $T_0 G V_1 G T_0 G V_1 G T_0$. The external amplitudes are only needed in half-off-shell kinematics, which comes out of the LO calculation. But the middle one is needed in kinematics where both the incoming and outgoing leg are off-shell, so that one deals (at a fixed on-shell point) with a pretty big matrix, and not with a vector. That was, I confess, just too much hassle for me.

To provide a bit more detail, let us construct another solution to eq. (2.3) in a slightly different way, namely by explicit but formal construction:

$$T = [\mathbb{1} - VG]^{-1}V \quad (2.23)$$

Inserting the expansions of T and V (and assuming for simplicity as before that $G \equiv G_0$), one finds

$$T_0 + \epsilon^1 T_1 + \epsilon^2 T_2 + \mathcal{O}(\epsilon^3) = \left[\mathbb{1} - V_0 G - \epsilon^1 V_1 G - \epsilon^2 V_2 G + \mathcal{O}(\epsilon^3) \right]^{-1} \left[V_0 + \epsilon^1 V_1 - \epsilon^2 V_2 + \mathcal{O}(\epsilon^3) \right] \quad (2.24)$$

One now expands both left- and right-hand side as before in powers of ϵ and keeps the inverse $[\mathbb{1} - V_0 G]^{-1}$ un-expanded. The LO is then the formal solution to the LO equation (??) we found above:

$$T_0 = [\mathbb{1} - V_0 G]^{-1} V_0 \quad (2.25)$$

Using this to make the equations at least a bit more compact, the corrections are

$$\begin{aligned}\mathcal{O}(\epsilon^1) : T_1 &= (\mathbb{1} + T_0 G) V_1 (\mathbb{1} + G T_0) \\ \mathcal{O}(\epsilon^2) : T_2 &= (\mathbb{1} + T_0 G) V_2 (\mathbb{1} + G T_0) + (\mathbb{1} + T_0 G) V_1 (\mathbb{1} + G T_0) G V_1 (\mathbb{1} + G T_0) \\ &\text{etc.}\end{aligned}\tag{2.26}$$

That is a lot of terms, but no amplitude beyond T_0 is involved on the right-hand side. You get the “double-off-shell” culprit $T_0 G V_1 G T_0 G V_1 G T_0$ as one of the contributions when you multiply out the last term on the right-hand side. When used in coordinate space, this approach is often identified as Distorted-Wave Born Approximation [19, sect. 9.1.2]

Fortunately, Jared Vanasse rediscovered the solution at N²LO in 2013 [16]. I and (I am sure) many colleagues as well as he himself quickly felt reminded of Mathematical Perturbation Theory and wrote down the general case at home. As in so many cases, that is to my knowledge not explicitly stated anywhere in the “modern” EFT literature, but it’s become something of a standard lore.

2.5 $ET = KT$: Eigenvalue Equations with Varying Eigenvalue

These are similar to eq. (2.46), but now the eigenvalue $E \neq 1$ can move and the parameter to be found is not hidden in T but the eigenvalue itself.

$$E T = K T \tag{2.27}$$

Implicitly, T and H are of course related to E . A prominent example is the time-independent Schrödinger equation, where $T(E)$ is the wave function and the kernel is (in contradistinction to the other sums, this one starts at $n = 1$!)

$$K = H_0 + V = H_0 + \epsilon^1 V_1 + \epsilon^2 V_2 + \mathcal{O}(\epsilon^3) = H_0 + \sum_{n=1} \epsilon^n V_n . \tag{2.28}$$

H_0 now stands for a LO interaction which contains both a potential V_0 and kinetic energy term, and V for all the potential which is not in H_0 – so here, $V = \epsilon V_1 + \mathcal{O}(\epsilon^2)$ since the LO piece of the interaction is already in H_0 . Likewise, we expand T as in eq. (2.6).

The continuum problem describes usually scattering at a fixed, order-independent eigenvalue (energy) E . In that case, the equations derived below for the “bound-state” case of an eigenvalue E which varies from order to order simplify since $E = E_0$ while $E_{n>0} = 0$.

Basing the inductive solution on an expansion of the LO eigensystem is discussed by any QM textbook, so I did originally not treat it in detail. These older notes of canonical QM perturbation are found in subsection 2.5.2. It is useful if the LO eigenvalues and eigenfunctions are all known (like for the harmonic oscillator or for the H atom), or if one is interested in systems with large gaps between the eigenvalue E_k of interest and other eigenvalues E_l , $|E_k - E_l| \gg |E_k|, |E_l|$. It often leads to analytic expressions which were particularly useful before computers became a research tool.

However, the systems of interest to more advanced research (nuclear, atomic, ...) do not have a closed-form LO solution, and the numerical problem favours a more-direct approach in sub-sect. 2.5.1 avoiding the construction of a complete ortho-normal system.

2.5.1 Solving Without Resort to a Complete Ortho-Normal System

Discussions with Sebastian König at an INT workshop in May 2026 alerted me to this more-efficient way. Triggered by my understanding of App. A and sect. II.D of Andis, Liu, Long and König [18], I present it based on Byron/Fuller [3, sec. 4.11] who provide a more pedagogical account but only cover the case of one perturbation, $V = V_0 + V_1$, albeit my generalisation below is not too difficult. I believe the article may imply a crucial step elaborated below but quite certainly does not actually describe it. This may be in part because it is not unlikely that the approach has been described and used before, but ref. [18] does not cite earlier cases and may well have assumed that this is “common knowledge”.

In this sub-section, we focus on a *particular eigenvalue* E of the full equation (2.27) and expand it as

$$E = E_0 + \epsilon^1 E_1 + \epsilon^2 E_2 + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n E_n . \quad (2.29)$$

Inserting this and eq. (2.28) into eq. (2.27), one finds

$$\begin{aligned} \mathcal{O}(\epsilon^0) : (H_0 - E_0)T_0 &= 0 \\ \mathcal{O}(\epsilon^1) : (H_0 - E_0)T_1 &= (E_1 - V_1)T_0 \\ \mathcal{O}(\epsilon^2) : (H_0 - E_0)T_2 &= (E_2 - V_2)T_0 + (E_1 - V_1)T_1 \\ \mathcal{O}(\epsilon^3) : (H_0 - E_0)T_3 &= (E_3 - V_3)T_0 + (E_2 - V_2)T_1 + (E_1 - V_1)T_2 \\ \mathcal{O}(\epsilon^n) : (H_0 - E_0)T_n &= (E_n - V_n)T_0 + \sum_{m=1}^{n-1} (E_m - V_m)T_{(n-m)} \end{aligned} \quad (2.30)$$

[I go here one order higher than usual for reasons that will become transparent later.] The left-hand side is identical at each order, constituting the homogeneous part of the equation and hence a traditional Schrödinger equation. T_n only appears on the left, so only LO has no inhomogeneous part; all others do and hence represent Schrödinger equations with “sources/driving terms” which are given by the lower-order solutions $T_{m < n}$ and energies $E_{m < n}$ and potentials $V_{m < n}$. So once one has a computational technique to solve a Schrödinger-like equation with inhomogeneity, constructing solutions is straightforward from the always identical inversion $[H_0 - E_0]^{-1}$, with an increasingly more involved inhomogeneous part.

Solving for Eigenvalues and Eigenvectors The inhomogeneity does also include the potential V_n of the order we seek, and, annoyingly, also the energy E_n . This seems to pose a bit of a conundrum because there is only one equation at N^{th} LO to solve for two unknowns E_n and T_n . The solution is already present in the QM textbook treatments, albeit it is often somewhat obscured by the use of a Complete Ortho-Normal System of the LO solution; see next subsection: Find E_n first, then T_n . As Byron/Fuller point out, this can be achieved by multiplying the n th-order equation from the left with the LO solution $T_0^\dagger$ and using the self-adjoint version of the LO equation, $T_0^\dagger(H_0 - E_0) = 0$, to set all left-hand terms to zero.

[This assumes of course that H_0 is self-adjoint, as was already assumed for K .]:

$$T_0^\dagger \sum_{m=1}^n (E_m - V_m) T_{(n-m)} = 0 \text{ for } n \neq 0 \quad . \quad (2.31)$$

[Notice change of upper summation limit for a compacter notation.] This works because, as Byron/Fuller [3, sect. 4.11, below eq. (4.55)] notice, all these equations are of the type $(H_0 - E_0)\vec{x} = \vec{y}$, with $\vec{y}$ the inhomogeneity (r.h.s. of eq. (2.30)), while E_0 is an eigenvalue to H_0 , $\det(H_0 - E_0) = 0$ for a nontrivial LO solution $T_0 \neq 0$. That is just the triviality that the LO solution must be orthogonal to any inhomogeneities.

Since the result is now independent of T_n , one directly solves for the eigenvalue corrections, having normalised the LO solution, $T_0^\dagger T_0 = 1$:

$$\begin{aligned} \mathcal{O}(\epsilon^1) : E_1 &= T_0^\dagger V_1 T_0 \\ \mathcal{O}(\epsilon^2) : E_2 &= T_0^\dagger \left[V_2 T_0 + (V_1 - E_1) T_1 \right] \\ \mathcal{O}(\epsilon^3) : E_3 &= T_0^\dagger \left[V_3 T_0 + (V_2 - E_2) T_1 + (V_1 - E_1) T_2 \right] \\ \mathcal{O}(\epsilon^n) : E_n &= T_0^\dagger \left[V_n T_0 + \sum_{m=1}^{n-1} (V_m - E_m) T_{(n-m)} \right] \end{aligned} \quad (2.32)$$

Readers will immediately recognise the first line as the result of old-fashioned first order perturbation in energy. As one is used from QM perturbation theory, the eigenvalue (energy) shifts E_n only depend on the power-order corrections $T_{m < n}$ and potentials $V_{m < n}$, plus a linear dependence on the N^n LO potential V_n . One now inserts these shifts back into eq. (2.30), so that one can solve for T_n as the only remaining unknown:

$$\begin{aligned} \mathcal{O}(\epsilon^1) : T_1 &= [H_0 - E_0]^{-1} (E_1 - V_1) T_0 \\ \mathcal{O}(\epsilon^2) : T_2 &= [H_0 - E_0]^{-1} \left[(E_2 - V_2) T_0 + (E_1 - V_1) T_1 \right] \\ \mathcal{O}(\epsilon^3) : T_2 &= [H_0 - E_0]^{-1} \left[(E_3 - V_3) T_0 + (E_2 - V_2) T_1 + (E_1 - V_1) T_2 \right] \\ \mathcal{O}(\epsilon^n) : T_n &= [H_0 - E_0]^{-1} \left[(E_n - V_n) T_0 + \sum_{m=1}^{n-1} (E_m - V_m) T_{(n-m)} \right] \end{aligned} \quad (2.33)$$

The $2n$ th and $(2n + 1)$ st Eigenvalues Depend only on Eigenvectors up to T_n
Byron/Fuller [3, below eq. (4.56)] discuss now a nice feature: The eigenvalue corrections E_3 in eq. (2.32) actually depend only on T_0 and T_1 , but not even on T_2 . So see that, we must only eliminate T_2 from the right-hand side, which we can do as long as all operators are

self-adjoint:

$$\begin{aligned}
T_0^\dagger(V_1 - E_1)T_2 &= \left[(V_1 - E_1)T_0\right]^\dagger T_2 \\
&= -\left[(H_0 - E_0)T_1\right]^\dagger T_2 \\
&= -T_1^\dagger\left[(H_0 - E_0)T_2\right] \\
&= +T_1^\dagger\left[(V_2 - E_2)T_0 + (V_1 - E_1)T_1\right]
\end{aligned} \tag{2.34}$$

where one uses eq. (2.30): Inserting its $\mathcal{O}(\epsilon^1)$ result into the second line produces the third line, and its $\mathcal{O}(\epsilon^2)$ result the last line. The 3rd-order correction is then

$$\begin{aligned}
E_3 &= T_0^\dagger\left[V_3T_0 + (V_2 - E_2)T_1\right] + T_1^\dagger\left[(V_2 - E_2)T_0 + (E_1 - V_1)T_1\right] \\
&= T_0^\dagger V_3 T_0 + 2\text{Re}[T_0^\dagger(V_2 - E_2)T_1] + T_1^\dagger(V_1 - E_1)T_1
\end{aligned} \tag{2.35}$$

where self-adjointness of V_2 (implying also that E_2 is real) is used in the last line. So, E_3 depends indeed only on the eN^3LO perturbation V_3 and previous perturbations and energies, but only on T_0 and T_1 , and not even on T_2 .

Byron/Fuller assert that this pattern extends to higher orders: For E_4 (or E_5), only T_0 to T_2 are needed together with the eigenvalues up to E_3 (or E_4) and perturbations up to V_4 (or V_5), and for E_{2n} and E_{2n+1} , only the eigenvectors T_0 to T_n are needed. Thus, the eigenvalue corrections are not as numerically sensitive to the lower-order solutions $T_{m < n}$ than one might fear, which substantially improves numerical stability. Extension of this pattern to any order seems more an exercise in bookkeeping than Physics, so it is skipped here.

This feature is not found in solving equations of the Lippmann-Schwinger type in sect. 2.2, but that does of course not necessarily imply that Schrödinger-like equations can be solved with higher accuracy/precision.

However, one is usually less interested in E_n than in its eigenvector-correction T_n since that translates into a wave function that is convoluted with other operators, like electromagnetic currents. Unfortunately I cannot see a similar trick to eliminate the dependence of T_3 on T_2 (and analogously for higher orders). But the reduced numerical errors in the determination of E_3 *etc.* also reduce numerical uncertainties in T_3 from eq. (2.33). While the standard lore is therefore correct that perturbation theory for the eigenvalue/energy is simple while perturbation theory for the eigenvector/wave function is a pain, if one is interested in eigenvalue/energy shifts, one can go to pretty high orders without too much of it. Of course, once one is interested in matrix elements like form factors or transition functions, one needs the eigenvectors/vertex functions/wave functions T_n and is hit by the proverbial “no pain, no gain” doctrine.

Ortho-Normalising Simplifies Eigenvalue Equations Further Byron/Fuller [3, sect. 4.11 below eq. (4.56)] describe another trick to simplify the eigenvalue equations. One can actu-

ally show that all eigenvector corrections must be orthogonal to the LO eigenvector:

$$T_0^\dagger T_n = 0 \text{ for all } n > 0 . \quad (2.36)$$

In textbooks, this trick is often referred as “any correction must be orthogonal to LO”. At first sight, we would expect that higher-order corrections contain components parallel to T_0 , but that is not true. The argument is quite simple: Let T_n be a solution to the n th line in eq. 2.30. Since the l.h.s. is $[H_0 - E_0]T_n$ and $[H_0 - T_0]T_0 = 0$ by construction, one can always choose a new $\tilde{T}_n = T_n - (T_0^\dagger T_n) T_0$ which is orthogonal to T_0 and still solves the n th equation. If one invokes such a system of perturbations T_n orthogonal to T_0 , then the eigenvalue equations of eq. (2.32) simplify because all terms $T_0^\dagger E_m T_n = 0$ disappear

$$\begin{aligned} \mathcal{O}(\epsilon^1) : E_1 &= T_0^\dagger V_1 T_0 \\ \mathcal{O}(\epsilon^2) : E_2 &= T_0^\dagger [V_2 T_0 + V_1 T_1] \\ \mathcal{O}(\epsilon^3) : E_3 &= T_0^\dagger [V_3 T_0 + V_2 T_1 + V_1 T_2] = T_0^\dagger V_3 T_0 + 2\text{Re}[T_0^\dagger V_2 T_1] + T_1^\dagger (V_1 - E_1) T_1 \\ \mathcal{O}(\epsilon^n) : E_n &= T_0^\dagger \left[V_n T_0 + \sum_{m=1}^{n-1} V_m T_{(n-m)} \right] \end{aligned} \quad (2.37)$$

subject to the constraint $T_0^\dagger T_n = 0$ for all $n > 0$.

[The $\mathcal{O}(\epsilon^3)$ line also uses eq. (2.35).] This does again not help in the determination of the T_n ’s, but it can help reducing numerical errors in finding the E_n ’s. Whether the cost of the auxiliary condition outweighs that benefit must be decided on the numerics of the specific problem at hand.

Normalising Eigenvectors Finally, the remark on normalising T_n in sect. 2.2.2 applies again, except that now we already imposed $T_0^\dagger T_0 = 1$ in eq. (2.32).

2.5.2 Solving Via A Complete Ortho-Normal System of Eigenfunctions to LO

For this, I find the account in Byron/Fuller [3, sects. 4.11 and 4.12] better than most QM treatments. It is also already pretty general, so the following is largely copied from there. [Byron/Fuller treat only a perturbation $V_1 = 0, V_{i>1} = 0$, but the extension below is a trivial generalisation.] I keep these notes largely in the 2018 version, more for historic than any other reason. So this part could be shortened, but that implies a re-write.

The equations of the preceding subsection provide a solution by matrix inversion and linear algebra, so why another method? Byron/Fuller [3, sect. 4.11 above eq. (4.28)] write: “The most obvious method would be to solve the matrix equations directly. However, since the main application of perturbation theory is to operators on infinite-dimensional spaces [their emphasis], where there is no simple matrix representation, we shall say no more about this direct method other than that it exists and is in principle easy to use on a finite-dimensional space.

If we were fortunate enough to know the complete set of eigenfunction and eigenvalues for the Hermitian [sic!] operator A_0 [our H_0], there would be another possibility:...” and continue to describe the approach reported now.

That was certainly an appropriate view at a time when computers were rare and slow, and the focus was on perturbations around analytic LO results. As pointed out above, the systems of interest to more advanced research (nuclear, atomic,...) do not have a closed-form LO solution, and the numerical problem favours a more-direct approach in sub-sect. 2.5.1 avoiding the construction of a complete ortho-normal system.

In the textbook treatment, one first constructs an eigensystem of $k = 1, \dots$ eigenvalues (energies) $\{E_{0k}\}$ and corresponding eigenvectors (wave functions) $\{T_{0k}\}$ to the LO problem involving H_0 for the k th state:

$$\mathcal{O}(\epsilon^0) : H_0 T_{0k} = E_{0k} T_{0k} \quad , \quad (2.38)$$

As the eigensystem changes from order to order, the k th eigenvalue of the full equation (2.27) is

$$E_k = E_{0k} + \epsilon^1 E_{1k} + \epsilon^2 E_{2k} + \mathcal{O}(\epsilon^3) = \sum_{n=0} \epsilon^n E_{nk} \quad . \quad (2.39)$$

The eigenvalues (and hence the index k) can be discrete or continuous, or both. The continuum problem describes usually scattering at a fixed, order-independent energy E ; *cf.* above. [If some of the eigenvalues are degenerate, you have fun with degenerate perturbation theory, including avoided level-crossing. I will not delve into that here; see [3, sect. 4.12].]

This eigensystem is then used to determine higher orders. At the higher orders, one finds after a few trivial rearrangements another pyramid that needs to be solved top-to-bottom:

$$\begin{aligned} \mathcal{O}(\epsilon^1) : (H_0 - E_{0k})T_{1k} &= (E_{1k} - V_1)T_{0k} \\ \mathcal{O}(\epsilon^2) : (H_0 - E_{0k})T_{2k} &= (E_{1k} - V_1)T_{1k} + (E_{2k} - V_2)T_{0k} \\ \mathcal{O}(\epsilon^n) : (H_0 - E_{0k})T_{nk} &= \sum_{m=1}^n (E_{mk} - V_m)T_{(n-m)k} \end{aligned} \quad (2.40)$$

Only the left-hand side contains the unknowns T_{nk} , and we apply the same LO eigenvalue operator $(H_0 - E_{0k})$ there at each order. However, the unknown E_{nk} appears on the right. How can we find these two unknowns at N^n LO from just one equation? Still, this is not even an integral equation, so we should be able to solve this. Byron/Fuller notice that all these equations are of the type $(H_0 - E_{0k})\vec{x} = \vec{y}$, so that, since E_{0k} are eigenvalues to H_0 , we must have $\det(H_0 - E_{0k}) = 0$ and $\vec{y}$ must be orthogonal to $\vec{x} = T_0$, while we still can expand it in the LO CONS $\vec{y} = \sum_k c_k T_{0k}$. In other symbols,

$$T_{0k}^\dagger T_{0k'} = \delta_{kk'} \text{ for normalised eigenvectors.} \quad (2.41)$$

We use this and eq. (2.38) to solve the problem with E_{nk} at N^n LO. We multiply each

equation in eq. (2.40) from the right with $T_{0k}^\dagger$ to find

$$\begin{aligned}\mathcal{O}(\epsilon^1) : E_{1k} &= T_{0k}^\dagger V_1 T_{0k} \\ \mathcal{O}(\epsilon^2) : E_{1k} &= T_{0k}^\dagger \left[V_2 T_{0k} + (V_1 - E_{1k}) T_{1k} \right] \\ \mathcal{O}(\epsilon^n) : E_{nk} &= T_{0k}^\dagger V_n T_{0k} + T_{0k}^\dagger \sum_{m=1}^{n-1} (V_m - E_{mk}) T_{(n-m)k}\end{aligned}\tag{2.42}$$

Now, T_{nk} disappeared at N^n LO. Notice the upper summation limit of $(n-1)$ in the sum.

To find the corrections T_{nk} to the eigenvectors, we recall that the set of all eigenvectors $\{E_{0k}\}$ to H_0 is a Complete Ortho-Normal Set, *i.e.* it spans the whole space (H_0 is self-adjoint since it's (usually) a QM operator related to an observable). One can therefore expand any vector into that basis:

$$T_{nk} = \sum_{k' \neq k} c_{nkk'} T_{0k'} \quad \text{with } c_{nkk'} = T_{0k'}^\dagger T_{nk}\tag{2.43}$$

We end up with a 3-index object $c_{nkk'}$ for the coefficient of the n th-order correction to the eigenvector k which points along the eigenvector $k' \neq k$ (the ordering of the indices matters...). Since eigenvectors to different eigenvalues are orthogonal, we write

$$T_{0k'}^\dagger (H_0 - E_{0k}) T_{0k} = (E_{0k'} - E_{0k}) T_{0k'}^\dagger T_{0k} = 0 \quad \text{for } k \neq k'\tag{2.44}$$

to find the coefficients $c_{nkk'}$ which allow us to reconstruct T_{nk} from eq. (2.43):

$$\begin{aligned}\mathcal{O}(\epsilon^1) : T_{1k} &= \sum_{k' \neq k} T_{0k'} \frac{1}{E_{0k} - E_{0k'}} T_{0k'}^\dagger V_1 T_{0k} \\ \mathcal{O}(\epsilon^n) : (E_{0k'} - E_{0k}) c_{nkk'} &= T_{0k'}^\dagger \sum_{m=1}^n (E_{mk} - V_m) T_{(n-m)k}\end{aligned}\tag{2.45}$$

which of course only works if the LO eigenvalues are not degenerate.

But what is more, for energy shifts up to order ϵ^3 , we do only need T_{0k} and T_{1k} ! One sees from eq. (2.42) that T_{0n} suffices to determine E_{0k} and E_{1k} . A close look at eq. (2.30) shows that adding only a solution T_{1n} fixes E_{2k} and E_{3k} . Indeed, each T_{nk} determines not one but 2 terms in the expansion of E_{nk} ; see the previous sub-section.

2.6 $T = KT$: Eigenvalue Equations with Fixed Eigenvalue

We now address the homogeneous part, $T = KT$ of the integral-type equations of sect. 2.2. These ‘‘Operator Equations of the Integral Type Without Inhomogeneity’’ only have solutions at some particular values of a parameter which is now called $\mathcal{E}$ in order to avoid confusion with the energy variable E :

$$T(\mathcal{E}) = K(\mathcal{E}) T(\mathcal{E}) \quad \text{i.e. } [\mathbb{1} - K(\mathcal{E})] T(\mathcal{E}) = 0 \quad \text{such that } T(\mathcal{E}) \neq 0 \text{ for some } \mathcal{E}\tag{2.46}$$

is an eigenvalue equation for a *kernel* K (GV in sect. 2.2) with eigenvalue 1 which remains unchanged at any order. The parameter $\mathcal{E}$ is adjusted order-by-order such that the equation contains at least one non-trivial solution. One example is the homogeneous equation related to eq. (2.3). Its solution is the vertex function T of a bound state whose energy is now indeed given by the parameter $\mathcal{E}$. Another example are Schrödinger-like equations at fixed energy but with a varying parameter $\mathcal{E}$ which enters (usually highly nonlinearly) in K ; the fixed energy can be set without loss of generality to $E = 1$; if not, rescale K .

Originally, discussion of the NLO correction copied with little change from ref. [17], but I am sure this had been done before as well. In July 2026, I realised that the generalisation to $N^n\text{LO}$ is actually quite straightforward, given the work in sects. 2.2 and 2.5.

The parameter $\mathcal{E}$ is tuned such that T is an eigenvector of K with eigenvalue 1. Following the Mathematical Perturbation Theory method, both $K(\mathcal{E})$, $T(\mathcal{E})$ and $\mathcal{E}$ are expanded in powers of ϵ . Let us assume the case that the matrix K contains perturbations both by itself and in the parameter:

$$\begin{aligned}
K(\mathcal{E}) &= K_0(\mathcal{E}) + \epsilon^1 K_1(\mathcal{E}) + \epsilon^2 K_2(\mathcal{E}) + \dots = K(\mathcal{E}_0 + \epsilon^1 \mathcal{E}_1 + \epsilon^2 \mathcal{E}_2 + \dots) \\
&= K_0 + \epsilon^1 [\mathcal{E}_1 K'_0 + K_1] + \epsilon^2 \left[\mathcal{E}_2 K'_0 + K_2 + \mathcal{E}_1 K'_1 + \frac{\mathcal{E}_1^2}{2} K''_0 \right] + \\
&\quad + \epsilon^3 \left[\mathcal{E}_3 K'_0 + K_3 + \mathcal{E}_2 K'_1 + \mathcal{E}_1 \left(K'_2 + \mathcal{E}_2 K''_0 + \frac{\mathcal{E}_1}{2!} K''_1 + \frac{\mathcal{E}_1^2}{3!} K^{(3)}_0 \right) \right] \dots \quad (2.47) \\
&=: \sum_{n=0} \epsilon^n \left[\mathcal{E}_n K'_0 + K_n + \mathcal{K}_n[(K_{m < n}^{(p \leq n)})^{l \leq n}; \mathcal{E}_{m < n}] \right] ,
\end{aligned}$$

where $K^{(n)} \equiv \frac{d^n}{d\mathcal{E}^n} K(\mathcal{E})|_{\mathcal{E}=\mathcal{E}_0}$ denotes the n th derivative of K at $\mathcal{E}_0$ and it is understood that each operator is evaluated at $\mathcal{E}_0$. [Notation: Each term inside the second and the last summation is again of order ϵ^n , as indicated by summing subscripts in multiplications.] The terms become quickly lengthy, so I am not writing down the general expansion. Instead, the last line groups terms into the correction of K at $N^n\text{LO}$, the corresponding correction to $\mathcal{E}$ at the same order, and combinations $\mathcal{K}_n$ of corrections which are individually of lower order but are multiplied to become n th order in total. The first few terms are read off se:

$$\begin{aligned}
\mathcal{O}(\epsilon^1) : \mathcal{K}_1 &= 0 \\
\mathcal{O}(\epsilon^2) : \mathcal{K}_2 &= \mathcal{E}_1 K'_1 + \frac{\mathcal{E}_1^2}{2} K''_0 \\
\mathcal{O}(\epsilon^3) : \mathcal{K}_3 &= \mathcal{E}_2 K'_1 + \mathcal{E}_1 \left(K'_2 + \mathcal{E}_2 K''_0 + \frac{\mathcal{E}_1}{2!} K''_1 + \frac{\mathcal{E}_1^2}{3!} K^{(3)}_0 \right)
\end{aligned} \quad (2.48)$$

We can now simply recycle previous results: The solution to (2.46) combines the variable-eigenvalue problem of eq. (2.30) with $H \rightarrow K$ and at fixed energy parameter $E_0 = 1$ and $E_{n>0} = 0$ with the nonlinear fitting problem for $K(\mathcal{E})$ of sect. 2.3.2 where one replaces

$V \rightarrow K(\mathcal{E})$ and expands K as in eq. (2.47), so that eq. (2.30) becomes:

$$\begin{aligned}
\mathcal{O}(\epsilon^0) : (K_0 - 1)T_0 &= 0 \quad \text{determining } \mathcal{E}_0 \\
\mathcal{O}(\epsilon^1) : (K_0 - 1)T_1 &= -(\mathcal{E}_1 K'_0 + K_1)T_0 \\
\mathcal{O}(\epsilon^2) : (K_0 - 1)T_2 &= -(\mathcal{E}_2 K'_0 + K_2 + \mathcal{K}_2)T_0 - (K_1 + \mathcal{E}_1 K'_0)T_1 \\
\mathcal{O}(\epsilon^n) : (K_0 - 1)T_n &= -(\mathcal{E}_n K'_0 + K_n + \mathcal{K}_n)T_0 + \sum_{m=1}^{n-1} \left(\mathcal{E}_m K_0^{(m)} + K_m + \mathcal{K}_m \right) T_{n-m}
\end{aligned} \tag{2.49}$$

One now yet again multiplies from the left with $T_0^\dagger$, uses the LO identity to set the l.h.s. to zero, and solves for the unknown. In the variable-eigenvalue problem, it was E_n ; now, we look for $\mathcal{E}_n$, with the LO equation simply giving $\mathcal{E}_0$:

$$\begin{aligned}
\mathcal{O}(\epsilon^1) : \mathcal{E}_1 &= -\frac{T_0^\dagger K_1 T_0}{T_0^\dagger K'_0 T_0} \\
\mathcal{O}(\epsilon^2) : \mathcal{E}_2 &= -\frac{T_0^\dagger [K_2 + \mathcal{K}_2] T_0 + T_0^\dagger [K_1 + \mathcal{E}_1 K'_0] T_1}{T_0^\dagger K'_0 T_0} \\
\mathcal{O}(\epsilon^n) : \mathcal{E}_n &= -\frac{T_0^\dagger [K_n + \mathcal{K}_n] T_0 + T_0^\dagger \left[\sum_{m=1}^{n-1} \left(\mathcal{E}_m K_0^{(m)} + K_m + \mathcal{K}_m \right) T_{n-m} \right]}{T_0^\dagger K'_0 T_0}
\end{aligned} \tag{2.50}$$

The NLO result agrees with what is reported in [17]. The denominator, $T_0^\dagger K'_0 T_0$, is identical at each order and can thus be computed once to high accuracy and stored. Since this is still the Schrödinger-like solution, each $\mathcal{E}_n$ only depends on the results from lower orders, and each T_n is found from eq. (2.49) after the corresponding $\mathcal{E}_n$ has been determined. Following the steps leading to eqs. (2.35), the results for $\mathcal{E}_2$ and $\mathcal{E}_3$ again depend only on T_0 and T_1 , or more generally those for $\mathcal{E}_{2n}$ and $\mathcal{E}_{2n+1}$ only on $T_{m \leq n}$, *e.g.*:

$$\mathcal{E}_3 = -\frac{T_0^\dagger [K_3 + \mathcal{K}_3] T_0 + 2\text{Re}[T_0^\dagger [\mathcal{E}_2 K'_0 + K_2 + \mathcal{K}_2] T_1] + T_1^\dagger [\mathcal{E}_1 K'_0 + K_1] T_1}{T_0^\dagger K'_0 T_0} \tag{2.51}$$

To reiterate like a broken record: If one is only interested in the shift of *e.g.* the binding energy of the homogeneous Lippmann-Schwinger equation, then finding $\mathcal{E}_n$ suffices and one needs a few $T_{[n/2]}$ on the way. If one is however interested in the eigenvectors T_n (*e.g.* as the vertex functions of the bound state which are related to wave functions), then one needs to resort to eq. (2.49).

2.7 Distorted Wave Born Approximation

This is very similar to the eigenvalue equation with varying eigenvalue; *cf.* [19, sect. 9.1.2]. It does not need an extra section.

2.8 More Variations

If you know of or are interested in other strictly perturbative expansions about nontrivial LO equations, let me know and I would be happy to add them – with full credit, of course.

3 From Solutions to Observables

As alluded to in sect. 2.2.2, one should take some care when constructing observables from solutions: Mathematical Perturbation Theory dictates that these be expanded consistently, with no order higher than the order of the solution. Including more terms does not increase the accuracy of the observable since other terms from interactions V_{n+1} *etc.* are absent. It is also simply inconsistent with the tenets of Mathematical Perturbation Theory. For now I refer to a recent article whose Appendix A details the order-by-order expansion of phase shifts [6], and only describe the procedure for moving poles.

When the solutions T represents the T -matrix/scattering amplitude, then I choose symbolically the following relation to the scattering/ S -matrix:

$$S = e^{2i\delta} = \mathbb{1} + 2i k T \text{ with } T = \frac{1}{k \cot \delta(k) - ik} \quad (3.1)$$

[Dropping the spurious factor $\frac{2\pi}{\mu}$ helps one to concentrate on the argument (μ is the reduced mass of the scattering problem).]

Below the first inelasticity, $k \cot \delta(k)$ is a real function of the scattering phase shift δ at centre-of-mass momentum k . This form assumes an uncoupled problem, but the extension for coupled channels is treated in the literature and not conceptually challenging; see *e.g.* [6–8] and references therein. Below the first inelasticity, $k \cot \delta(k)$ is a real function of the scattering phase shift δ at centre-of-mass momentum k .

3.1 Unitarity

It is sometimes claimed that a perturbative treatment “violates unitarity”. That is both correct and wrong – it may better be said that the statement is irrelevant. Yes, the perturbative treatment will not exactly fulfil $|T|^2 = 1$ (for wave functions) or $|\mathbb{1} + T|^2 = 1$ (for T -matrices), but that does not lead to imaginary parts for example in phase shifts *if one follows the mandate of Mathematical Perturbation Theory that all variables must be expanded and all orders must be kept consistently for all observables*. Indeed, we discuss now that this provides a good check of numerical accuracy.

We now have a solution $T = T_0 + \epsilon^1 T_1 + \epsilon^2 T_2 + \dots$ of the left-hand side. Mathematical Perturbation Theory requires to also expand the right-hand side and match orders in ϵ . For a direct test of unitarity, one may test if the interaction part of the amplitude is indeed real, treating “ $k \cot \delta$ ” as one variable and expanding it:

$$(k \cot \delta) = (k \cot \delta)_0 + \epsilon^1 (k \cot \delta)_1 + \epsilon^2 (k \cot \delta)_2 + \dots \quad (3.2)$$

Invert eq. (3.1) and expand,

$$\begin{aligned} (k \cot \delta)_0 + \epsilon^1 (k \cot \delta)_1 + \epsilon^2 (k \cot \delta)_2 + \dots &= ik + \frac{1}{T_0 + \epsilon^1 T_1 + \epsilon^2 T_2 + \dots} \\ &= ik + \frac{1}{T_0} - \epsilon^1 \frac{T_1}{T_0^2} + \epsilon^2 \left[\frac{T_1^2}{T_0^3} - \frac{T_2}{T_0^2} \right] + \dots \end{aligned} \quad (3.3)$$

so that

$$\begin{aligned} (k \cot \delta)_0 &= ik + \frac{1}{T_0} \\ (k \cot \delta)_1 &= -\frac{T_1}{T_0^2} \\ (k \cot \delta)_2 &= \frac{T_1^2}{T_0^3} - \frac{T_2}{T_0^2} \end{aligned} \quad (3.4)$$

[Notation: The order of each multiplication/division continues to be given by the sum of its subscripts.] The LO result is by construction exactly unitary, $ik + \frac{1}{T_0} \in \mathbb{R}$. The second term is the NLO correction, the third one the N²LO one. Constructed this way, all higher-order corrections *must* be real, with a imaginary part which is a numerical zero. This provides a stringent test of both numerics and correctness of implementation.

As an exercise, consider Bethe's Effective Range Expansion: $k \cot \delta = -\frac{1}{a} + \frac{r}{2} k^2 + \dots$, where the effective range r is a NLO correction to LO, which in turn is parametrised by the scattering length a . One convinces oneself quickly that as expected

$$\begin{aligned} (k \cot \delta)_0 &= -\frac{1}{a} \\ (k \cot \delta)_1 &= \frac{r}{2} k^2, \end{aligned} \quad (3.5)$$

i.e. the interaction function is real, as is the phase shift. Another test of unitarity is of course to check if the phase shift δ is real below the first inelasticity.

3.2 Extracting Phase Shifts

One has choice how to perturbatively extract phase shifts. Different choices of an expansion up to $\mathcal{O}(\epsilon^n)$ must agree up to terms of $\mathcal{O}(\epsilon^{n+1})$; if not, then assumptions of Perturbation Theory does not hold; see sect. 2.1.

For example, one can simply solve the expansion of the interaction function,

$$\delta = \operatorname{arccot} \frac{(k \cot \delta)_0}{k} + \operatorname{arccot} \frac{(k \cot \delta)_1}{k} + \operatorname{arccot} \frac{(k \cot \delta)_2}{k} + \dots, \quad (3.6)$$

where each term is one order. Alternatively, one can expand directly

$$\delta = \delta_0 + \epsilon^1 \delta_1 + \epsilon^2 \delta_2 + \dots \quad (3.7)$$

and insert that into the relation (3.1) between S - and T -matrix, so that [subscript S for extraction using S -matrix]:

$$\begin{aligned}\delta_{0S} &= -\frac{i}{2} \ln [1 + 2i k T_0] \\ \delta_{1S} &= \frac{k T_1}{1 + 2i k T_0} = k T_1 e^{-2i\delta_0} \\ \delta_{2S} &= \frac{k T_2}{1 + 2i k T_0} - i (\delta_1)^2 = k T_2 e^{-2i\delta_0} - i (\delta_1)^2 .\end{aligned}\tag{3.8}$$

The denominator is always a power of $S = 1 + 2i k T_0 = e^{2i\delta_0}$. One can test again perturbative unitarity via the degree to which the δ_n are real.

All expansion variants must be identical at LO: Since $S_0 = 1 + 2i k T_0$ is unitary at LO, there is no ambiguity in $\delta_0 \bmod \pi$: $\delta_0 = \delta_{0S}$. The reader is invited to convince oneself that the NLO differences are compensated at N²LO; and N²LO differences will be cured when one includes N³LO terms.

Details and formulae for coupled channels are found in *e.g.* [6–8] and references therein.

3.3 Finding and Moving Poles in Amplitudes

3.3.1 Method

In QFT, bound and resonance states are characterised by the poles and cuts in the S -matrix, or, equivalently, in the amplitude T . For simplicity, I assume an on-shell amplitude between two particles, and that it depends on one momentum k ; extension to N -point amplitudes/Green’s functions is...you guessed it: “left as an exercise to the reader”. This presentation expand on handwritten notes I had prepared for some participants of the Nanjing workshop in 2018. As I started writing these notes, I was sure that encountered the following method first in the context of the KSW articles and hallway conversations in Seattle. However, I was as of yet unable to find a citeable source, but I am not claiming what I write is new. It would be truly embarrassing if the first document were indeed [20].

While it is often said that poles cannot be shifted in perturbation theory, commonplace QM contradicts that. Atomic levels are poles in the atom’s S -matrix. Calculating the Stark and Zeeman effects of energy shifts of atomic levels in external fields provides thus an obvious example that poles can shift in perturbation theory; see methods 2.6 and 2.5.

One identifies poles as zeroes of the inverse scattering amplitude. Let the pole be located at momentum $k = i\gamma$, where $\gamma > 0$ for a real bound state of binding energy $B = \gamma^2/M$ ($B > 0$ when bound), $\gamma < 0$ for a virtual bound state, and $\gamma \in \mathbb{C}$ for a resonance with width $\text{Im}[\gamma]$. Let

$$T = T_{-1} + \epsilon T_0 + \epsilon^2 T_1 + \mathcal{O}(\epsilon^3)\tag{3.9}$$

be the amplitude, where T_{-1} is LO and $\epsilon^2 T_1 \ll \epsilon T_0 \ll T_{-1}$ the N²LO and NLO corrections, respectively, in an EFT power-counting scheme with expansion parameter ϵ . The extension to higher orders is trivial bookkeeping of which I do not want to deprive the reader as it will no doubt be enjoyed beyond measure.

Following EFT procedure, an expansion exists for inverse amplitude, like for any observable, in powers of the expansion parameter. [Yes, an amplitude is not quite an observable, but it comes very near one. In particular, the positions of poles and cuts are physical observables.] [The expression involves operator/matrix inversions as needed.]

$$\frac{1}{T_{-1} + \epsilon T_0 + \mathcal{O}(\epsilon^2)} = \frac{1}{T_{-1}} - \epsilon \frac{T_0}{(T_{-1})^2} + \epsilon^2 \frac{T_0^2 - T_{-1}T_1}{(T_{-1})^3} + \mathcal{O}(\epsilon^3) \quad (3.10)$$

One now identifies a pole at $k = i\gamma$ by requiring that this function is zero. As before, it may be advantageous to expand the pole position $\gamma = \gamma_{-1} + \epsilon \gamma_0 + \epsilon^2 \gamma_1 + \mathcal{O}(\epsilon^3)$, insert into eq. (3.10) and expand the inverse amplitude, for example:

$$0 \stackrel{!}{=} \frac{1}{T_{-1}(\gamma_{-1} + \epsilon \gamma_0)} = \frac{1}{T_{-1}(\gamma_{-1})} - \epsilon \gamma_0 \left. \frac{d}{dk} \frac{1}{T_{-1}(k)} \right|_{k=\gamma_{-1}} + \mathcal{O}(\epsilon^2) \text{ etc.} \quad (3.11)$$

The result must then be zero at each order of ϵ independently:

$$\begin{aligned} \mathcal{O}(\epsilon^0) : 0 &\stackrel{!}{=} \frac{1}{T_{-1}(\gamma_{-1})} \\ \mathcal{O}(\epsilon^1) : 0 &\stackrel{!}{=} -\gamma_0 \left. \frac{d}{dk} \frac{1}{T_{-1}(k)} \right|_{k=\gamma_{-1}} - \frac{T_0(\gamma_{-1})}{(T_{-1}(\gamma_{-1}))^2} \implies \gamma_0 = -\frac{T_0(\gamma_{-1})}{T'_{-1}(\gamma_{-1})} \\ &\text{etc.} \end{aligned} \quad (3.12)$$

All amplitudes and derivatives are to be taken at γ_{-1} , and primes denote derivatives with respect to k . As before, the first equation determines γ_{-1} , which is then inserted into the second one to determine γ_0 , and so forth. Notice that we take derivatives of the inverse of T at pole positions. These are well-defined, while one might have mathematical headaches to take derivatives of analytic functions right at their pole. In that sense, expressions like T'_{-1} involved in the solution for γ_0 should be taken with a grain of salt.

3.3.2 Example: NLO Effective Range Theory

Warning: I use in this sub-section a T -matrix with *opposite sign* from eq. (3.1), for historical reasons. I should re-write this to be consistent. For now, flip every $T \rightarrow -T$.

To illuminate a few points, consider non-relativistic scattering between identical particles. The amplitude up to NLO in EFT($\not{\Lambda}$) is found from the effective-range expansion.

$$\begin{aligned} T(k) &= \frac{1}{\gamma + ik - \frac{r}{2} k^2} = \frac{1}{\gamma + ik} + \frac{\frac{r}{2} k^2}{(\gamma + ik)^2} + \mathcal{O}\left(\left(\frac{r}{a}\right)^2\right) \\ \implies T_{-1}(k) &= \frac{1}{\gamma + ik}, \quad T_0(k) = \frac{\frac{r}{2} k^2}{(\gamma + ik)^2} = \frac{r}{2} k^2 T_{-1}^2(k) \end{aligned} \quad (3.13)$$

where γ is traditionally interpreted as inverse scattering length and r as effective range, setting the breakdown scale of the theory as $\frac{1}{r} \sim \Lambda_{\not{\Lambda}} \gg k$.

Both T_{-1} and T_0 have a pole at $k = i\gamma$. So naïvely, there is no way that the pole can be shifted at NLO. However, it should make one suspicious that T_0 has actually a second-order pole, and on a moment's notice one realises that any higher-order amplitude T_n will have poles of even higher order $n + 2 \geq 2$. This is clearly wrong. Physical poles must be of first order. The fact that Nature allows physical particles tells us that the argument is wrong, unless perturbation theory itself is wrong.

It is not too difficult to see the root of the problem. By the EFT tenet, higher-order corrections must be parametrically suppressed. The ratio

$$\frac{T_0(k)}{T_{-1}(k)} = \frac{rk}{2} \frac{1}{\frac{\gamma}{k} + i} \quad (3.14)$$

is indeed smaller than 1 in magnitude for $|rk|/2 < 1$ when $k \lesssim \gamma$ – *unless* $k \rightarrow i\gamma$, in which case the second term explodes. Consider an expansion $k = i\gamma + dk$ about the pole. One then has $|T_{-1}| > |T_0|$ only outside the circle about $k = i\gamma$ with radius $|dk| = |\frac{r}{2} \gamma^2|$. Stated differently, a second-order pole diverges faster than any first-order pole sufficiently close to the singularity: $T_0 \gg T_{-1}$ close to the pole, $|dk| \leq |\frac{r}{2} \gamma^2|$, in an apparent violation of the EFT power-counting¹. However, the expression

$$\frac{1}{T_{-1} + T_0} \rightarrow \frac{1}{T_{-1}(k)} - \frac{T_0}{(T_{-1})^2} + \mathcal{O}(\epsilon^2) = (\gamma + ik) + \frac{r}{2} k^2 + \mathcal{O}(\epsilon^2) \quad (3.15)$$

is perfectly well-behaved and perturbative at the pole, as long as $|rk|/2 < 1$, i.e. k is below the breakdown scale $\Lambda_\pi \sim 1/r$. The alleged second-order pole in T_0 is exactly cancelled by the double-zero of $1/T_{-1}^2$.

The LO pole position is now found from eq. (3.10) at LO,

$$0 \stackrel{!}{=} \gamma + ik \implies k_0^{\text{LO}} = i\gamma, \quad (3.16)$$

and the NLO pole position from eq. (3.10) at NLO,

$$0 \stackrel{!}{=} (\gamma + ik) + \frac{r}{2} k^2 \implies k_0^{\text{NLO}} = i\gamma \left[1 - \frac{\gamma r}{2} \right]. \quad (3.17)$$

At NLO, the pole has shifted from its LO position. This is of course nothing but the good old relation between the deuteron binding energy and the effective-range parameters:

$$B = \frac{\gamma^2}{M} \rightarrow \frac{\gamma^2}{M} [1 + \gamma r + \mathcal{O}((\gamma r)^2)] \quad (3.18)$$

[Exercise: Use the perturbation formalism of eq. (3.12) to find the same result.]

¹In response to this argument, Ubirajara van Kolck tells me, if I understand him correctly, that he interprets my procedure as in effect resumming the perturbations close to the pole in order to arrive at eq. (3.10). I think that is a fair interpretation which does however not affect the mathematical technique. One can think of the eq. (3.10) as taking the analytic continuation of the amplitude expansion into the circle around the pole where eq. (3.9) does not converge.

In this particular example, the pole position at NLO is perturbatively close to the LO position γ for $|\gamma r| \ll 1$, and the LO result γ actually sets the scale of the shift. Since $|\gamma r| \ll 1$, the pole position may change, but a real bound state will remain real, and a virtual one will remain virtual.

That, however, is not always the case. For example, one may start from the unitarity limit $\gamma = 0$ at LO. Then, the N²LO perturbations restore the actual finite scattering length, $\gamma \neq 0$ and an effective range, moving the LO pole right from threshold $B = 0$ to the LO effective-range result $B = \frac{\gamma^2}{M}$; see [21]. It is then the sign of the perturbative correction $\gamma \neq 0$ which determines if the bound state is real or virtual. The corresponding LO and NLO amplitudes are

$$T_{-1} = \frac{1}{ik} \quad , \quad T_0 = \frac{1}{(ik)^2} \left(\frac{r}{2} k^2 - \gamma \right) = \left(\frac{r}{2} k^2 - \gamma \right) T_{-1}^2 \quad . \quad (3.19)$$

We can devise a very simple *model* which moves a virtual bound state to a real one. Take as one amplitude

$$T_a = \frac{1}{\gamma_a + i k} \quad (3.20)$$

with some parameter γ_a , exhibiting a pole at $k_a = i\gamma_a$. Let another amplitude be

$$T_b = \frac{1}{\gamma_b + i k} + \frac{\frac{r}{2} k^2}{(\gamma_b + i k)^2} \quad (3.21)$$

with two parameters γ_b and r and exhibiting thus a pole at $k_b = i\gamma_b(1 - \frac{\gamma_b r}{2})$. T_a is a LO amplitude, T_b an NLO amplitude. We are now provided one datum at some momentum $k_{\text{fit}} > 0$, for example the value and slope of $k \cot \delta$. We use that to fix the parameters. The two pieces of information determine γ_b, r ; in T_a , we can only accommodate one of the experimental constraints, and we choose the value of the scattering amplitude. So at the datum k_{fit} , the two amplitudes are by construction identical:

$$-\gamma_a \stackrel{!}{=} -\gamma_b + \frac{r k_{\text{fit}}^2}{2} \quad (3.22)$$

The fit point k_{fit} is real for $\gamma_b > \gamma_a$ and $r > 0$ (or both < 0).

We can now choose a model in which $\gamma_b > 0$, *i.e.* we find a real bound state from amplitude T_b , while $\gamma_a < 0$, *i.e.* the bound state is virtual from amplitude T_a . In this model, therefore, the same phase shift information at some arbitrary but given point k_{fit} leads to not just different pole positions, but these positions change from virtual to real, *i.e.* are on different Riemann sheets (I had to throw that in somewhere...). We can smoothly interpolate between the two scenarios by defining an amplitude

$$T(t) = \frac{1}{\gamma_a + t(\gamma_b - \gamma_a) + i k} + \frac{\frac{r}{2} k^2 t}{(\gamma_a + t(\gamma_b - \gamma_a) + i k)^2} \quad (3.23)$$

for which $T(t=0) = T_a$ and $T(t=1) = T_b$ and the pole moves with $t \in [0; 1]$ as

$$k_{\text{pole}}(t) = i[\gamma_a + t(\gamma_b - \gamma_a)] \left(1 - [\gamma_a + t(\gamma_b - \gamma_a)] \frac{rt}{2} \right) \quad (3.24)$$

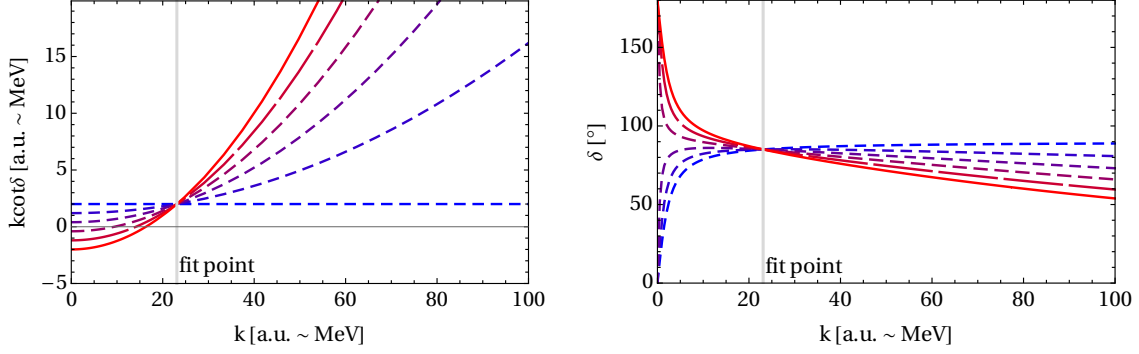


Figure 1: (Colour online) $k \cot \delta$ (left) and phase shift (right) as function of momentum k in the model of eq. (3.23) from $T(t=0) = T_a$ (blue) to $T(t=1) = T_b$ (red) in steps of $\Delta t = 0.2$. The parameters are set to $\gamma_a = -2$, $\gamma_b = 2$ and $r = \frac{3}{200}$ in some arbitrary units whose base resembles MeV.

from the negative imaginary axis at $t=0$ to the positive one at $t=1$, and exact unitarity at $t_U = \frac{\gamma_a - \gamma_b}{\gamma_a}$. If we had information at $k=0$ in this model, then, of course, $\gamma_a = \gamma_b$ and a real (virtual) bound state would remain real (virtual) under any higher-order correction. For any given t , the NLO amplitude provides a parametrically small correction to LO.

Alternatively, one can expand the effective-range about the pole at γ_a ,

$$\begin{aligned} \frac{1}{\gamma_a + \Delta\gamma + ik - \frac{r}{2} k^2} &\rightarrow \frac{1}{\gamma_a + ik} + \frac{\frac{r}{2} k^2 - \Delta\gamma}{(\gamma_a + ik)^2} \\ \Rightarrow T_{-1} &= \frac{1}{\gamma_a + ik} \quad , \quad T_0 = \left(\frac{r}{2} k^2 - \Delta\gamma \right) T_{-1}^2 \end{aligned} \quad (3.25)$$

Inserting this into eq. (3.10) at LO gives

$$0 \stackrel{!}{=} \gamma_a + i k \quad \Rightarrow \quad k_{\text{pole}}^{\text{LO}} = i\gamma_a \quad , \quad (3.26)$$

namely a real/virtual bound state for $\gamma_a > 0/\gamma_a < 0$. On the other hand, at NLO

$$0 \stackrel{!}{=} \gamma_a + i k + \frac{r}{2} k^2 - \Delta\gamma \quad \Rightarrow \quad k_{\text{pole}}^{\text{NLO}} = i(\gamma_a - \Delta\gamma) \left[1 - \frac{(\gamma_a - \Delta\gamma)r}{2} \right] \quad . \quad (3.27)$$

One sees again that the pole position k_{pole} can change between real ($\gamma_a - \Delta\gamma > 0$) and virtual ($\gamma_a - \Delta\gamma < 0$). But this now *only* depends on the relative size of γ_a to the correction $\Delta\gamma$ to the LO scattering length. In particular, we can still have $\gamma_a, \Delta\gamma < \frac{1}{r} \sim \Lambda_{\pi}$, *i.e.* both parameters are low-energy variables.

I feel like I am overlooking something really simple here – I am sure the bee hive will tell me what is crappy and why this model does not apply.

3.3.3 Example: Zeeman and Paschen-Back Effect in Hydrogen

Literally a textbook example of moving states from bound to the continuum is the hydrogen atom in a uniform magnetic field. The following discussion heavily draws from [22]. In the Zeeman limit, the field is weak and moves the binding energies of states only by small amounts. In the Paschen-Back limit, the field is so strong that states with $m_l + 2m_s < 0$ will eventually acquire sufficient energy to unbind. For example, the $n = 2$ state has binding energy $\frac{13.6}{2^2}$ eV = 3.4 eV at zero magnetic field, but the energy of its p -wave components with $m_l = -1, m_s = -\frac{1}{2}$ moves into the continuum; see fig. 5.3 of [22].

4 Summary and Conclusions

We did great things!

Acknowledgements

These notes grew out of a number of hallway discussion at many conferences, but I am particularly grateful to the organisers and participants of the workshop NEW IDEAS IN CONSTRAINING NUCLEAR FORCES at the ECT* in Trento in 2018 and of the workshop ESNT WORKSHOP TOWARDS HIGH-PRECISION IN MULTI-CHANNEL REACTIONS OF COMPOSITE NUCLEI at CEA-Saclay in 2026. If these notes are useful, it is because they insisted on them and on clarity. If they are useless, they are so due to my faults. Ragnar Stroberg spotted numerous initial typos and unclear points on that draft. The section on moving poles in amplitudes came out of notes I had prepared prompted by and circulated for discussions with the organisers and participants of the workshop EFFECTIVE FIELD THEORIES AND AB-INITIO CALCULATIONS OF NUCLEI at Nanjing University in 2019. Ubirajara van Kolck and Johannes Kirscher provided extensive comments on the scanned notes. Sebastian König pointed to the solution of Schrödinger-like equations without resort to a LO CONS at the INT PROGRAM INT-26-1 NUCLEAR HAMILTONIANS FOR ADVANCING NUCLEAR PHYSICS AND BEYOND in Seattle in 2026. My blackboard presentation at the workshop ESNT WORKSHOP TOWARDS HIGH-PRECISION IN MULTI-CHANNEL REACTIONS OF COMPOSITE NUCLEI at CEA-Saclay in 2026 was likely more understandable and definitely more focused than these notes. Comments by Jaume Carbonell and others prompted me to explore more avenues and augment these notes. This work was supported in part by the US Department of Energy under contract DE-SC0015393, and by the Dean's Research Chair programme of the Columbian College of Arts and Sciences of The George Washington University. It was in part conducted in GW's Campus in the Closet.

A Future Appendices

References

- [1] C. M. Bender and S. A. Orszag: Advanced Mathematica Methods for Scientists and Engineers.
- [2] J. M. Murdock: Perturbations: Theory and Methods.
- [3] F. W. Byron and R. W. Fuller: Mathematics of Classical and Quantum Physics.
- [4] H. W. Griebhammer: PHYS 6110: Graduate Mathematical Methods of Theoretical Physics, Chapter I, George Washington University; manu-script at <http://home.gwu.edu/~hgrie>
- [5] D. J. Griffiths and D. F. Schroeter: *Introduction to Quantum Mechanics*, 3rd ed., Cambridge University Press 2018.
- [6] Y. P. Teng and H. W. Griebhammer, Eur. Phys. J. A **61** (2025) no.9, 211 doi:[10.1140/epja/s10050-025-01675-6](https://doi.org/10.1140/epja/s10050-025-01675-6) [[arXiv:2410.09653](https://arxiv.org/abs/2410.09653) [nucl-th]].
- [7] O. Thim, A. Ekström and C. Forssén, Phys. Rev. C **109** (2024) no.6, 064001 doi:[10.1103/PhysRevC.109.064001](https://doi.org/10.1103/PhysRevC.109.064001) [[arXiv:2402.15325](https://arxiv.org/abs/2402.15325) [nucl-th]].
- [8] O. Thim, “*Chiral Effective Field Theory with Partly Perturbative Pions Applied to the Few-Nucleon Sector*,” PhD thesis, Chalmers U., doi:[10.63959/chalmers.dt/5891](https://doi.org/10.63959/chalmers.dt/5891)
- [9] H. W. Griebhammer, J. A. McGovern and D. R. Phillips, Eur. Phys. J. A **54** (2018) 37 doi:[10.1140/epja/i2018-12467-8](https://doi.org/10.1140/epja/i2018-12467-8) [[arXiv:1711.11546](https://arxiv.org/abs/1711.11546) [nucl-th]].
- [10] V. Pascalutsa and D. R. Phillips, Phys. Rev. C **67** (2003) 055202 [[nucl-th/0212024](https://arxiv.org/abs/nucl-th/0212024)].
- [11] H. W. Griebhammer, J. A. McGovern, D. R. Phillips and G. Feldman, Prog. Part. Nucl. Phys. **67** (2012) 841 [[arXiv:1203.6834](https://arxiv.org/abs/1203.6834) [nucl-th]].
- [12] J. A. McGovern, D. R. Phillips and H. W. Griebhammer, Eur. Phys. J. A **49** (2013) 12 doi:[10.1140/epja/i2013-13012-1](https://doi.org/10.1140/epja/i2013-13012-1) [[arXiv:1210.4104](https://arxiv.org/abs/1210.4104) [nucl-th]].
- [13] H. W. Griebhammer, **PoS(CD15)** 104 [[arXiv:1511.00490](https://arxiv.org/abs/1511.00490) [nucl-th]].
- [14] T. Mehen and I. W. Stewart, Phys. Lett. B **445** (1999) 378 doi:[10.1016/S0370-2693\(98\)01470-1](https://doi.org/10.1016/S0370-2693(98)01470-1) [[nucl-th/9809071](https://arxiv.org/abs/nucl-th/9809071)].
- [15] P. F. Bedaque, G. Rupak, H. W. Griebhammer and H. W. Hammer, Nucl. Phys. A **714** (2003) 589 doi:[10.1016/S0375-9474\(02\)01402-1](https://doi.org/10.1016/S0375-9474(02)01402-1) [[nucl-th/0207034](https://arxiv.org/abs/nucl-th/0207034)].
- [16] J. Vanasse, Phys. Rev. C **88** (2013) no.4, 044001 doi:[10.1103/PhysRevC.88.044001](https://doi.org/10.1103/PhysRevC.88.044001) [[arXiv:1305.0283](https://arxiv.org/abs/1305.0283) [nucl-th]].

- [17] S. Kreuzer and H. W. Griesshammer, Eur. Phys. J. A **48** (2012) 93
doi:[10.1140/epja/i2012-12093-6](https://doi.org/10.1140/epja/i2012-12093-6) [[arXiv:1205.0277](https://arxiv.org/abs/1205.0277) [nucl-th]].
- [18] A. J. Andis, S. Lyu, B. Long and S. König, [[arXiv:2512.12823](https://arxiv.org/abs/2512.12823) [nucl-th]].
- [19] R. G. Newton: Scattering Theory of Waves and Particles, 2nd ed.
- [20] P. F. Bedaque and H. W. Griesshammer, Nucl. Phys. A **671**, 357 (2000)
doi:[10.1016/S0375-9474\(99\)00691-0](https://doi.org/10.1016/S0375-9474(99)00691-0) [nucl-th/9907077].
- [21] S. König, H. W. Griesshammer, H. W. Hammer and U. van Kolck, Phys. Rev. Lett. **118**,
no. 20, 202501 (2017) doi:[10.1103/PhysRevLett.118.202501](https://doi.org/10.1103/PhysRevLett.118.202501) [[arXiv:1607.04623](https://arxiv.org/abs/1607.04623) [nucl-
th]].
- [22] J. J. Sakurai and J. Napolitano: Modern Quantum Mechanics, 3rd ed.