

Re-iterative Methods For Integral Equations

Paul Antony Burton

This dissertation is submitted in partial
fulfilment of the requirements of
the degree of Master of Science.

Department of Mathematics.

University of Reading.

September, 1993.

Abstract

A review is given of the Galerkin, iterated Galerkin and re-iterated Galerkin methods for finding approximate solutions to integral equations of the second kind. The process of re-iteration, first described by Porter & Stirling, is applied to the Kantorovich method to produce two new re-iteration methods. The first, the modified iterated Kantorovich method, is a straightforward application of the re-iterated Galerkin method to a regularized Kantorovich equation. The second method differs by applying the Kantorovich regularization at each iteration. This second method is tested on some example integral equations and its convergence compared with that of the re-iterated Galerkin method. Finally it is shown how the presence of an eigenvalue can affect the convergence of these methods.

Acknowledgements

I wish to thank, first and foremost, my supervisor for this project, Dr.D.Porter, for his invaluable help and encouragement throughout. I would also like to thank my fellow MSc students and members of the department who have assisted me over the last year. Finally I wish to acknowledge the financial support of the Science and Engineering Research Council.

Contents

Table of Notation	iv
1 <u>Introduction</u>	1
1.1 Integral equations	1
1.2 Projection methods	3
2 <u>Galerkin Methods</u>	4
2.1 The Galerkin Method	4
2.2 Degenerate Kernels	5
2.3 The Iterated Galerkin Method	7
2.4 The Re-iterated Galerkin Method	8
3 <u>Kantorovich Methods</u>	13
3.1 The Kantorovich Method	13
3.2 The Iterated Kantorovich Method	15
3.3 The Modified Iterated Kantorovich Method	16
3.4 The Re-iterated Kantorovich Method	19
4 <u>Numerical Considerations</u>	26
4.1 Numerical Integration	26

4.2	Choosing The Subspace	27
5	<u>Examples and Results</u>	29
5.1	The Kernel $2\max(x,t)$	29
5.2	The $\frac{1}{2} \sin(\kappa_0 x-t)$ Kernel	33
6	<u>Convergence Regions and Eigenvalue Bounds</u>	37
6.1	The Kernel $2\lambda\max(x,t)$	38
6.1.1	Calculating $\rho(S)$	39
6.2	The Kernel $\frac{\lambda}{2} \sin(\kappa_0 x-t)$	44
7	<u>Extending to 2-D</u>	47
8	<u>Conclusions</u>	51

Table of Notation

$k(x, t)$	kernel
K	integral operator
S, SK	convergence operators
μ_1^+, μ_1^-	eigenvalues
$\tilde{\mu}_1$	approximate eigenvalue
(ϕ, ψ)	inner product
$\ \cdot\ $	norm
H	Hilbert space
$L_2(0, 1)$	space of square integrable functions
$R_K(\phi)$	Rayleigh quotient
I	identity operator
P_n	projection
$\rho(T)$	spectral radius of T
$\kappa_0, \alpha, \beta, \lambda$	parameters
$\phi(x)$	particular function
E_n	n-dimensional subspace
$\{\chi_1, \dots, \chi_n\}$	basis functions
sup	supremum

In certain problems it is not uncommon to encounter equations in which the un-

b

a

b

a

1

0

We wish to find approximations to the solution of integral equations of the form

$$= \quad + \quad (1\ 3)$$

in some Hilbert space , such as ${}_2(0\ 1)$.

To do this we look for an approximation u_n to u in some finite dimensional subspace V_n of V , by solving the equation

$$n = n + n \quad (14)$$

in $\mathcal{H}_n$, where P_n is an orthogonal projection of $\mathcal{H}$ onto $\mathcal{H}_n$.

Chapter 2

Galerkin Methods

2.1 The Galerkin Method

If our approximation p_n is chosen such that $Ap_n - f$ is orthogonal to the subspace E_n , then we have

$$\begin{aligned}(Ap_n - f, \chi_j) &= 0, \quad j = 1, \dots, n \\ \Rightarrow (Ap_n, \chi_j) &= (f, \chi_j), \quad j = 1, \dots, n\end{aligned}$$

where the χ_j 's are the orthonormal basis functions of E_n .

Using our definition of p_n from equation (1.5) and $A = I - K$ (taking $\lambda = 1$ in equation (1.3) for convenience) we arrive at

$$\sum_{i=1}^n \alpha_i \{(\chi_i, \chi_j) - (K\chi_i, \chi_j)\} = (f, \chi_j), \quad j = 1, \dots, n \quad (2.1)$$

This system of n equations can be solved simultaneously, provided the determinant of coefficients is non-zero, to obtain the α_i 's which, via equation (1.5), provide us with our approximation p_n to the solution ϕ of our integral equation.

In terms of the projection P_n of H onto E_n ,

$$P_n p_n = P_n f + P_n K p_n$$

$$n = \left(\begin{matrix} & n \\ n & \end{matrix} \right)^{-1} n$$

$$\text{since } n \cdot n = n.$$

$$/$$

$$n$$

$$n$$

$$N \qquad N \ N$$

$$N$$

$$\sum_{i=1}^N \quad i \quad i$$

$$N$$

$$N$$

Notice that this equation is identical to the expression for the $\hat{u}_n$'s in Galerkin's method, so the $\hat{u}_n$'s are the same for both methods. However, $\hat{u}_N = \hat{u}_N$. In fact

$$\hat{u}_N = \hat{u}_0 + \sum_{n=1}^N \hat{u}_n = \hat{u}_0 + \hat{u}_N \quad (2.5)$$

$$\frac{n}{n-1}$$

$$n$$

$$\begin{array}{cccccccccccc} 0 & & 0 & 0 & & 1 & & 0 & 1 & 0 & 1 \\ & & & 0 & & & & 0 & & 1 & \\ & & & & n & 1 & & 0 & & & n & 1 & 0 \\ & & & & & & & & & & & & \\ & & n & & n & 1 & & & n & & n & 1 & 0 \end{array}$$

$$\frac{1}{0} \qquad n \qquad n \qquad 1$$

$$n \qquad n \qquad 0 \qquad n \qquad 0$$

$$\frac{\hat{n}}{\hat{n}-1}$$

in the initial (non-iterated) Galerkin approximations we find, (taking $n = 0$),

$$\begin{aligned}
r_0 &= f + Kp_0 - p_0 \\
&= \hat{p}_0 - p_0 \\
&= (I - KP)^{-1}f - (I - P_nK)^{-1}P_nK \\
&= (I - P_n)(I - KP_n)^{-1}f,
\end{aligned}$$

which leads us to

$$\frac{\|\hat{r}_n\|}{\|r_n\|} \leq \|K\|, \tag{2.9}$$

ie. the ratio of the residual of the iterated Galerkin approximation, over the norm of the residual of the Galerkin approximation, provides a lower bound on the norm of the operator K . Whereas this quantity has no direct consequences for the convergence of the method, it is however valuable information to know about the kernel, and comes relatively free.

The re-iterated Galerkin method provides a means of calculating an approximation to the solution of an integral equation, providing certain conditions are met, to any desired accuracy. The number of re-iterations required to reach a specific accuracy depends upon the rate of convergence of the method, which, in turn, depends upon the size of the subspace used. An increase in the size of $\mathcal{V}_n$ will result in faster convergence, but requires more calculations.

Is there a way in which the re-iteration process can be made more efficient, increasing the rate of convergence whilst minimising the number of calculations?

equation (3.1) represents the application of the operator K to both sides of the original integral equation $\phi = f + K\phi$.

We get a clue that we may have improved matters by considering the Neumann series for our equation, given by

$$\phi = f + Kf + K^2f + \dots + K^n f + \dots \quad (3.2)$$

The Kantorovich method gives $\phi = f + Kf + K\psi$, compared to $\phi = f + K\phi$ for Galerkin, and so contains an extra term of the Neumann series.

We approximate ψ by finding the Galerkin solution of the equation $\psi = Kf + K\psi$, ie.

$$\psi \approx q_n = \sum_{i=1}^n \beta_i \chi_i,$$

where $q_n \in E_n$, a subspace spanned by the orthonormal basis $\chi_1, \dots, \chi_n$.

We want to choose q_n so that $A\psi - Kf$ is orthogonal to E_n , where $A = I - K$. Therefore

$$(Aq_n - Kf, \chi_j) = 0, \quad j = 1, \dots, n$$

$$(Aq_n, \chi_j) = (Kf, \chi_j), \quad j = 1, \dots, n$$

$$\sum_{i=1}^n \beta_i (\chi_i, \chi_j) - (K\chi_i, \chi_j) = (Kf, \chi_j), \quad j = 1, \dots, n$$

Notice that the β_i 's given by the system of equations above differ from the β_i 's given by Galerkin's method because we have a different free term in the integral equation, (Kf, χ_j) instead of (f, χ_j) . The solution to our original problem is therefore approximated by $\tilde{\phi}_n = f + q_n$, which is different from the ϕ_n of Galerkin's method.

The benefit of making the substitution in the Kantorovich method becomes clear when we compare the differences between the approximations and $\phi - \phi_n$.

$(I - K)^{-1}f$. If we denote by P_n the orthogonal projection onto E_n , we see that

$$p_n - \phi = (I - P_n K)^{-1}(P_n - I)(I - K)^{-1}f, \quad (3.3)$$

and

$$\tilde{p}_n - \phi = (I - P_n K)^{-1}(P_n K - K)(I - K)^{-1}f. \quad (3.4)$$

The difference between the two errors is that for the first, P_n has to be chosen such that $(P_n - I)(I - K)^{-1}f$ is small, so the choice depends on f , whereas, for the Kantorovich method, it is only necessary to choose P_n such that $\|P_n K - K\|$ is small, independently of f , which is more easily satisfied (see Sloan [6]).

Another advantage of the Kantorovich method is that it removes f from the calculations, replacing it by Kf . This is desirable if f is not a particularly smooth function. The application of K to f has the effect of smoothing the free term. Kf is then used in the calculations to approximate ψ , with f being added on afterwards to give $\tilde{p}$, the approximation to ϕ .

3.2 The Iterated Kantorovich Method

Sloan [6] demonstrated that the principle of iteration, which he first applied to Galerkin's method, could also be brought to bear on Kantorovich's method.

We wish to find the solution $\phi = f + \psi$, where ψ satisfies the integral equation $\psi = Kf + K\psi$. The Galerkin approximation to ψ gives

$$\psi \approx q_0 = (I - P_n K)^{-1}P_n Kf \quad \Rightarrow \quad \phi \approx \tilde{p}_0 = f + q_0.$$

If we apply the iteration process to q_0 , by replacing ψ in the equation, we arrive at

$$\hat{q}_0 = Kf + Kq_0$$

ie.

$$\begin{aligned}\hat{q}_0 &= (I - KP_n)^{-1}Kf \\ \Rightarrow \quad \phi &\approx \hat{p}_0 = f + \hat{q}_0.\end{aligned}$$

The error in the approximation of ψ is given by

$$\begin{aligned}\hat{q}_0 - \psi &= (I - KP_n)^{-1}Kf - (I - K)^{-1}Kf \\ &= (I - KP_n)^{-1}(K - KP_n)(q_0 - \psi) \\ &= S(q_0 - \psi), \quad \text{where } S = (I - KP_n)^{-1}(K - KP_n),\end{aligned}$$

which in turn implies

$$\hat{p}_0 - \phi = S(p_0 - \phi).$$

The operator S is the same as that encountered in the iterated Galerkin method, which is as one might expect considering that the two methods differ only in the choice of free term. As with the Galerkin method, iteration will improve the Kantorovich approximation if the norm of the operator S is less than one, which can be guaranteed by a suitable choice of subspace. The advantage of the iterated Kantorovich method over the iterated Galerkin method lies in the possibility that $\psi = \phi - f$ will be smoother than ϕ , and thus easier to approximate.

3.3 The Modified Iterated Kantorovich Method

Although it has not previously been demonstrated, it seems fair to suppose that, since Kantorovich's method benefits from iteration, it might equally benefit from the process of re-iteration.

The iterated Kantorovich method is obtained by applying the iterated Galerkin method to the equation $\psi = Kf + K\psi$, where $\psi = \phi - f$. The iterated approximation $\hat{q}_0$ exhibits an error which can be compensated for by adding a correction term, ψ_1 say, so that

$$\psi = \hat{q}_0 + \psi_1.$$

By substituting ψ into the modified integral equation we see

$$\begin{aligned}\psi_1 &= Kf + K\hat{q}_0 - \hat{q}_0 + K\psi_1 \\ &= \hat{r}_0 + K\psi_1,\end{aligned}$$

where $\hat{r}_0$ is the residual error incurred in approximating ψ by $\hat{q}_0$.

We now have an integral equation to solve for ψ_1 , with free term

$$\begin{aligned}\hat{r}_0 &= Kf - (I - K)\hat{q}_0 \\ &= K(Kf - (I - K)q_0) \\ &= Kr_0.\end{aligned}$$

As before, the ratio of the residual norms of the iterated and non-iterated approximations estimate $\|K\|$.

Once again we can use iterated Galerkin to find an approximation $\hat{q}_1$ to ϕ_1 , and hence our new approximated solution to the modified equation becomes

$$\psi \approx \hat{q}_0 + \hat{q}_1,$$

where $\hat{q}_1 = (I - KP_n)^{-1}\hat{r}_0$, which has an error

$$\begin{aligned}\hat{q}_0 + \hat{q}_1 - \psi &= (I - KP_n)^{-1}\hat{r}_0 + (\hat{q}_0 - \psi) \\ &= -(I - KP_n)^{-1}(I - K)(\hat{q}_0 - \psi) + (\hat{q}_0 - \psi)\end{aligned}$$

$$\begin{aligned}
&= (\hat{A}_n)^{-1}(\hat{A}_n)(\hat{A}_0) \\
&= (\hat{A}_0) \\
&= \hat{A}_0^2(\hat{A}_0)
\end{aligned}$$

ie. $\hat{A}_1 = \hat{A}_0^2(\hat{A}_0)$, where $\hat{A}_1 = \hat{A}_0 + \hat{A}_1$. So $\hat{A}_0 + \hat{A}_1$ is an improved approximation to A if $\|A - (\hat{A}_0 + \hat{A}_1)\| < 1$, which we have previously assumed.

As before, this process can be repeated by noting that

$$A - (\hat{A}_0 + \hat{A}_1) = \hat{A}_2$$

where $\hat{A}_2$ is a correction term which satisfies the equation

$$\hat{A}_2 = \hat{A}_1 + \hat{A}_2$$

where $\hat{A}_1 = \hat{A}_0 (\hat{A}_0)^{-1} \hat{A}_1 = (\hat{A}_n)^{-1} \hat{A}_0$, the residual error in $\hat{A}_0 + \hat{A}_1$. As with the re-iterated Galerkin method, the ratio of successive residual norms, $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$, gives an underestimate for the norm of $\hat{A}_n$.

In general, then, we have $\hat{A}_{kant}^n = \hat{A}_0 + \hat{A}_1 + \dots + \hat{A}_n$, where each correction term $\hat{A}_n$ is determined by using iterated Galerkin on $\hat{A}_n = \hat{A}_{n-1} + \hat{A}_n$, giving $\hat{A}_n = (\hat{A}_n)^{-1} \hat{A}_{n-1}$, where $\hat{A}_{n-1} = (\hat{A}_{kant}^{n-1})$. The error in the current approximation being given by $\hat{A}_{kant}^n - A = (\hat{A}_{kant}^{n-1}) = \hat{A}_{n+1}(\hat{A}_0)$.

The benefits of re-iteration for the Kantorovich method are the same as for the

generally smaller than that for the re-iterated Galerkin method, $\epsilon_0 \approx 10^{-4}$. All subsequent approximations, however, improve by the same factor, $\approx 10^{-4}$, for both methods. So the rates of convergence are virtually the same.

The improvement in the initial approximation, brought about by the Kantorovich "regularization" of the original equation, leads us to wonder whether subsequent approximations might also benefit from similar regularization.

Suppose, once again, that we wish to solve

$$u = Au + f \quad (3.5)$$

in a Hilbert space H , where A is a compact linear map on H , $f \in H$ is given, and u is sought. We first perform a Kantorovich "regularization" on the above equation, whereby we write the solution as

$$u = u_0 + v \quad (3.6)$$

where $u_0 = A^{-1}f$ satisfies the equation

$$u_0 = Au_0 + f \quad (3.7)$$

We seek a Galerkin approximation $\tilde{u}_0$ to u_0 in a subspace H_n of H , so that $\tilde{u}_0 = (P_n A^{-1} P_n)^{-1} P_n f$, where P_n is the orthogonal projection of H onto H_n . Our Kantorovich approximation, $\tilde{u}_0$, is therefore

$$\tilde{u}_0 = \tilde{u}_0 + v$$

Next we form the iterated Galerkin solution of equation (3.7), ie.

$$\hat{u}_0 = (P_n A^{-1} P_n)^{-1} P_n f$$

giving us the iterated Kantorovich approximation

$$\phi \approx \hat{\hat{p}}_0 = \hat{\hat{q}}_0 + f,$$

where

$$\hat{\hat{p}}_0 - \phi = S(\hat{p}_0 - \phi).$$

As we have seen before, iteration improves the initial approximation by a factor $\|S\|$ for both the Kantorovich and Galerkin methods, the Kantorovich method having the possible advantage that $\hat{p}_0$ has better convergence properties than p_0 , since it relies upon $\|P_n K - K\| \rightarrow 0$ as $n \rightarrow \infty$, rather than $(P_n - I)(I - K)^{-1} \rightarrow 0$ as $n \rightarrow \infty$.

We can utilise this convergence property in a re-iterative method by performing a Kantorovich regularization at each iteration. Instead of finding a correction term ψ_1 to $\hat{\hat{q}}_0$, to form $\psi = \hat{\hat{q}}_0 + \psi_1$, as we did for the modified iterated Kantorovich method, we now look for a correction term for ϕ .

We know $\phi \approx \hat{\hat{p}}_0 = f + \hat{\hat{q}}_0$, so this implies that

$$\phi = \hat{\hat{p}}_0 + \phi_1, \tag{3.8}$$

which in turn implies that

$$\phi_1 = \hat{r}_0 + K\phi_1, \tag{3.9}$$

where $\hat{r}_0 = f - (I - K)\hat{\hat{p}}_0 = Kr_0$.

This method differs from the modified iterated Kantorovich method in that we now perform a second Kantorovich regularization, this time to equation (3.9).

Let $K\phi_1 = \psi_1$, then $\phi_1 = \hat{r}_0 + \psi_1$, where ψ_1 satisfies

$$\psi_1 = K\hat{r}_0 + K\psi_1.$$

The Galerkin solution to this is

$$_1 \hat{u}_1 = (\quad_n \quad)^{-1} \quad_n \hat{u}_0$$

which we can iterate to obtain

$$_1 \hat{u}_1 = \hat{u}_0 + \quad_1 = (\quad_n)^{-1} \hat{u}_0$$

The correction term $\quad_1$ is therefore given by

$$_1 \hat{u}_1 = \hat{u}_0 + \hat{u}_1$$

$$_1$$

$$_0 \quad_1$$

$$\begin{array}{ccccccc} _0 & _1 & & _0 & _1 & & _0 \\ & & & _0 & & & _{-1} \quad_0 \\ & & & & _n & ^{-1} & _0 \\ & & & & _n & ^{-1} & _n \\ & & & & & _2 & ^{-1} \quad_0 \\ & & & & _n & ^{-1} & \\ & & & & & _n & _0 \\ & & & & & & _0 \end{array}$$

$$_0 \quad_1 \quad_0 \quad_0$$

$$_0 \quad_1 \quad_2$$

where ϕ_2 is a correction term which satisfies the equation

$$\phi_2 = \hat{r}_1 + K\phi_2, \quad (3.11)$$

where $\hat{\phi}_1 = (\hat{\phi}_0 + \hat{\phi}_1) = \hat{\phi}_1$. Performing a Kantorovich regularization on the above equation gives us

$$\phi_2 = \hat{\phi}_1 + \phi_2$$

where ϕ_2 satisfies

$$\phi_2 = \hat{\phi}_1 + \phi_2$$

Applying Galerkin's method gives us $\phi_2 \approx \tilde{\phi}_2 = (\phi_n)^{-1} \phi_n \hat{\phi}_1$, and upon iteration, $\phi_2 \approx \hat{\phi}_2 = \hat{\phi}_1 + \tilde{\phi}_2 = (\phi_n)^{-1} \hat{\phi}_1$.

So now

$$\phi_2 \approx \hat{\phi}_2 = \hat{\phi}_1 + \hat{\phi}_2$$

which implies that

$$\hat{\phi}_0 + \hat{\phi}_1 + \hat{\phi}_2$$

The error in this approximation can be shown to be

$$\begin{aligned} \hat{\phi}_0 + \hat{\phi}_1 + \hat{\phi}_2 &= (\phi_n)^{-1} (\phi_n) (\hat{\phi}_0 + \hat{\phi}_1) \\ &= (\hat{\phi}_0 + \hat{\phi}_1) \end{aligned}$$

so that

$$\hat{\phi}_0 + \hat{\phi}_1 + \hat{\phi}_2 \approx \hat{\phi}_0 + \hat{\phi}_1$$

provided, once again, that $\phi_n \approx 1$.

In general, if $\hat{\phi}_{rekant}^n = \hat{\phi}_0 + \dots + \hat{\phi}_n$, and $\hat{\phi}_n = (\phi_n)^{\hat{\phi}_{rekant}^n}$ denotes the residual after the nth iteration, then the correction term $\phi_n = \hat{\phi}_{n-1} + \phi_n$, is

approximated by $\hat{u}_n = \hat{u}_{n-1} + \hat{r}_n$, where $\hat{r}_n = (I - P_n)^{-1} \hat{r}_{n-1}$ is the iterated solution of the regularized equation $P_n \hat{r}_n = \hat{r}_{n-1} + \hat{r}_n$.

This gives

$$\begin{aligned} \hat{u}_{n+1}^{rekant} &= \hat{u}_n^{rekant} \\ &= \hat{u}_0^{n+1} + \hat{r}_0^{n+1} \end{aligned}$$

Note that the operator controlling the convergence of this method is $(I - P_n)^{-1}$, rather than $(I - P_n)$ produced by the re-iterated Galerkin and modified iterated Kantorovich methods. This gives us a new condition on the size of the subspace V_n . We now require that V_n is chosen such that $\| (I - P_n)^{-1} \| \leq 1$.

As commented on previously, these norms give a weaker bound on the con-

$$\frac{\| \hat{u}_n \|}{\| \hat{u}_{n-1} \|} \leq \frac{\| \hat{u}_n \|}{\| \hat{u}_{n-1} \|} \leq 1$$

$$\begin{aligned} \hat{u}_n^{gal} &= \hat{u}_0^{gal} + \hat{r}_0^{gal} \\ \hat{u}_n^{kant} &= \hat{u}_0^{kant} + \hat{r}_0^{kant} \\ \hat{u}_n^{rekant} &= \hat{u}_0^{rekant} + \hat{r}_0^{rekant} \end{aligned}$$

Then, for $i=0$,

$$\hat{\sigma}_{gal}^0 = f + K p_0,$$

$$\hat{\sigma}_{kant}^0 = f + K f + K q_0,$$

$$\hat{\sigma}_{rekant}^0 = f + K f + K \tilde{q}_0,$$

and for $i=1$,

$$\hat{\sigma}_{gal}^1 = f + K f + K^2 p_0 + K p_1,$$

$$\hat{\sigma}_{kant}^1 = f + K f + K^2 f + K^2 q_0 + K q_1,$$

$$\hat{\sigma}_{rekant}^1 = f + K f + K^2 f + K^3 f + K^3 \tilde{q}_0 + K \tilde{q}_1.$$

In general, for $i=n$,

$$\begin{aligned}\hat{\sigma}_{gal}^n &= \sum_{j=0}^n K^j f + \sum_{j=0}^n K_{j+1} p_{n-j}, \\ \hat{\sigma}_{kant}^n &= \sum_{j=0}^{n+1} K^j f + \sum_{j=0}^n K_{j+1} q_{n-j}, \\ \hat{\sigma}_{rekant}^n &= \sum_{j=0}^{2n+1} K^j f + \sum_{j=0}^n K^{2j+1} \tilde{q}_{n-j}.\end{aligned}$$

From this we can see that after n iterations, the re-iterated Galerkin approximation contains the first $n+1$ terms of the Neumann series for $\phi = f + K\phi$, the modified iterated Kantorovich approximation contains the first $n+2$ terms, and the re-iterated Kantorovich approximation contains the first $2n+2$ terms of the Neumann series, twice as many as for the re-iterated Galerkin method. This property appears to suggest that, under suitable conditions, the re-iterated Kantorovich method will have a much faster rate of convergence than the other two methods. In reality, if $\|K\|$ is greater than one, the Neumann series diverges, and it is the other terms which have to compensate. The Neumann series itself cannot therefore be responsible for the rate of convergence.

In the following chapters we will be comparing the efficiency and accuracy of the re-iterated Galerkin method with the new re-iterated Kantorovich method. The modified iterated Kantorovich method has not been considered further due to the belief that any improvement over the re-iterated Galerkin method will be minimal.

Chapter

Numerical Considerations

4.1 Numerical Integration

The process of calculating successive corrections to previous approximations, a feature of all the re-iterative methods, ensures that the accuracy of the numerical solution can be increased until it is of the same order as the computational accuracy of the method. Once the residual $\hat{r}_n$ reaches this accuracy there is no point in continuing the iteration, as no further improvement can be made. The main factor governing the computational accuracy is the method of numerical integration used. In all the cases described here modified Gauss-Legendre quadrature was employed, using 10 points and 20 subintervals. In this way, the value of every function used in the method is known at each of the two hundred Gauss points in the interval.

Numerical integration routines tend to encounter difficulties when dealing with kernels containing slope discontinuities, such as

$$(Kf)(x) = \int_0^1 \max(x, t)f(t)dt, \quad 0 \leq x \leq 1,$$

which has such a discontinuity at $\alpha = 0$.

The problem may be remedied by noting that

$$\begin{aligned} \langle \psi | \hat{H} | \psi \rangle &= \int_0^1 \langle \psi | \hat{H} | \psi \rangle d\alpha \\ &= \int_0^1 \langle \psi | \hat{H} | \psi \rangle d\alpha + \int_0^1 \langle \psi | \hat{H} | \psi \rangle d\alpha \end{aligned}$$

This reformulated equation has a first term with a continuous first derivative, and a second term that may be evaluated exactly. (See Chamberlain [1], who used this device in the case where $\langle \psi | \hat{H} | \psi \rangle = \frac{1}{2\alpha} \langle \psi | \hat{H} | \psi \rangle$).

In order to keep calculations as simple as possible, attention is restricted to finding approximations in subspaces of minimal dimension. Provided that the subspace is of sufficient size to ensure the conditions $\langle \psi | \hat{H} | \psi \rangle \leq 1$ and $\langle \psi | \hat{H} | \psi \rangle \geq 1$ are satisfied,

$$\sum_{n=1}^{\infty} \frac{1}{n^2}$$

$$\sum_{n=1}^{\infty} \frac{1}{n^2} = \frac{\pi^2}{6}$$

would be a good approximation to ϕ . For this reason we conclude that we ought to be looking for approximations belonging to the n -dimensional subspace spanned by $\phi^2, \dots, \phi^{n-1}$. Indeed, it has been shown (Porter & Stirling [3]) that this constitutes the optimal choice of subspace.

For our one-dimensional case, therefore, we are looking for Galerkin approximations of the form $\phi = \phi_1 + \phi_2$, where ϕ_1 is chosen to be the free term of the integral equation. For the re-iterated Galerkin method this is obviously $\phi_1 = \phi$. What is perhaps not so obvious is what we mean by the free term for the Kantorovich method.

Whilst it is true we are trying to find ϕ satisfying $\phi = \phi_1 + \phi_2$, which would suggest taking $\phi_1 = \phi$, the Kantorovich regularization effectively removes ϕ_1 from the approximation, so that we are instead seeking to approximate ϕ_2 , where ϕ_2 satisfies $\phi_2 = \phi - \phi_1$. Here the free term is not ϕ , but ϕ_2 , so we should perhaps take our trial function to be $\phi_2 = \phi - \phi_1$ instead.

Since the aim of this project is to give a comparison of the different methods, it seems sensible that they should use the same subspaces. On the other hand, each

	0	1	2	3	4	$\frac{\ \hat{r}_n\ }{\ \hat{r}_{n-1}\ }$	$\frac{\ \hat{r}_n\ }{\ r_n\ }$
$\hat{r}_{rekant}^0$	2 500001	2 627627	3 032719	3 785712	4 996088		0 356
$\hat{r}_{rekant}^2$	1 813364	1 927510	2 283396	2 926562	3 932512	0 179	0 354
$\hat{r}_{rekant}^4$	1 791510	1 905215	2 259509	2 899160	3 898587	0 179	0 354
$\hat{r}_{rekant}^6$	1 790809	1 904501	2 258744	2 898282	3 897500	0 179	0 354
$\hat{r}_{rekant}^8$	1 790787	1 904478	2 258719	2 898254	3 897465	0 179	0 354
$\hat{r}_{rekant}^9$	1 790786	1 904477	2 258718	2 898253	3 897464	0 179	0 354
()	1 790786	1 904477	2 258718	2 898253	3 897464		

Table 5.2: Re-iterated Kantorovich approximations (=)

$$\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$$

$$\frac{\|\hat{r}_n\|}{\|r_n\|}$$

	0	1	2	3	4	$\frac{\ \hat{r}_n\ }{\ \hat{r}_{n-1}\ }$	$\frac{\ \hat{r}_n\ }{\ r_n\ }$
$\hat{r}_{rekant}^0$							
$\hat{r}_{rekant}^2$							
$\hat{r}_{rekant}^4$							
$\hat{r}_{rekant}^6$							

the projection P_n , and thus the choice of subspace, is approximated as 0.058. This is smaller still than is obtained by taking $\chi = f$. As we might expect, therefore, the approximations to ϕ , given by $\hat{\sigma}_{rekant}^n = \hat{p}_0 + \cdots + \hat{p}_n$, converge to an accuracy of six decimal places after only seven iterations. (Naturally the rate of convergence for all three methods would be improved by taking a higher dimensional subspace.)

The value given here for $\|K\|$ is noticeably different from that given in the two previous tables. It isn't immediately obvious why this should be; K hasn't changed, so the reason must be linked to the choice of subspace. In order to see what is happening we need to look at the Rayleigh quotient, denoted $R(\phi)$.

If T is a bounded, self-adjoint linear map from a Hilbert space, H , to itself, then the Rayleigh quotient

$$R_T(\phi) = \frac{(T\phi, \phi)}{(\phi, \phi)},$$

gives a lower bound on the largest positive eigenvalue of T , denoted $\mu_1^+(T)$. That is

$$\begin{aligned} \mu_1^+(T) &= \sup\{R_K(\phi) : \phi \neq 0\} \\ &= \max\{R_K(\phi) : \phi \neq 0\} : \mu_1^+(T) \neq 0. \end{aligned}$$

Also

$$\|T\| = \max\{|R_K(\phi)| : \phi \neq 0\}.$$

(See Porter – Stirling [3] Lemma 5.1)

If we take $\phi = r_n$, and $T = K^2$, we see that

$$R_K(r_n) = \frac{(Kr_n, Kr_n)}{(r_n, r_n)} \leq \|K\|^2 = \mu_1^+(K^2),$$

where φ_1^+ denotes the largest positive eigenvalue of $\mathcal{A}_n$, with φ_1^+ being the corresponding eigenvector. This implies that

$$\frac{\| \hat{r}_n \|^2}{\| r_n \|^2} = \frac{\| \hat{r}_n \|^2}{\| r_n \|^2} \leq \frac{\varphi_1^+}{\varphi_1^+} = 1$$

$$\frac{\| \hat{r}_n \|^2}{\| r_n \|^2} \leq \frac{\varphi_1^+}{\varphi_1^+} = 1$$

$$n$$

$$n \leq \frac{\varphi_1^+}{\varphi_1^+} = 1$$

$$\frac{1}{0}$$

$$\frac{1}{0} \leq \frac{1}{0} = 0$$

$$0$$

$$i \leq i$$

where

$$\begin{array}{rcc} & & 0 \\ i & & 0 \\ & 1 & \\ 0 & \hline & 0 & 0 \end{array}$$

$$0$$

$$\begin{array}{ccccccc} & & // & & 2 & & \\ & & & & 0 & & \\ 0 & & 0 & 0 & 0 & ' & 0 & 0 \\ \iota(1) & & 0 & 0 & 0 & ' & 0 & 0 \end{array}$$

$$\begin{array}{cccc} \hline 0 & & 0 & 0 & 0 \\ \hline 0 & & 0 & 0 & 0 \\ \hline 2 & & 0 & & \\ 0 & & & & \end{array}$$

$$\begin{array}{ccccccc} 2 & 2 & & & & & \\ 0 & & 0 & & 0 & & 0 \end{array}$$

						_____ 1	_____

						_____	_____

The results in the previous chapter lead us to wonder in what cases the Kantorovich based re-iterative methods are better, in the sense of convergence rates,

$$\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$$

6.1 The kernel $2\lambda \max(\mathbf{x}, \mathbf{t})$

The first equation we consider is

$$\phi(x) = 1 + \lambda \int_0^1 2 \max(x, t) \phi(t) dt, \quad 0 \leq x \leq 1. \quad (6.1)$$

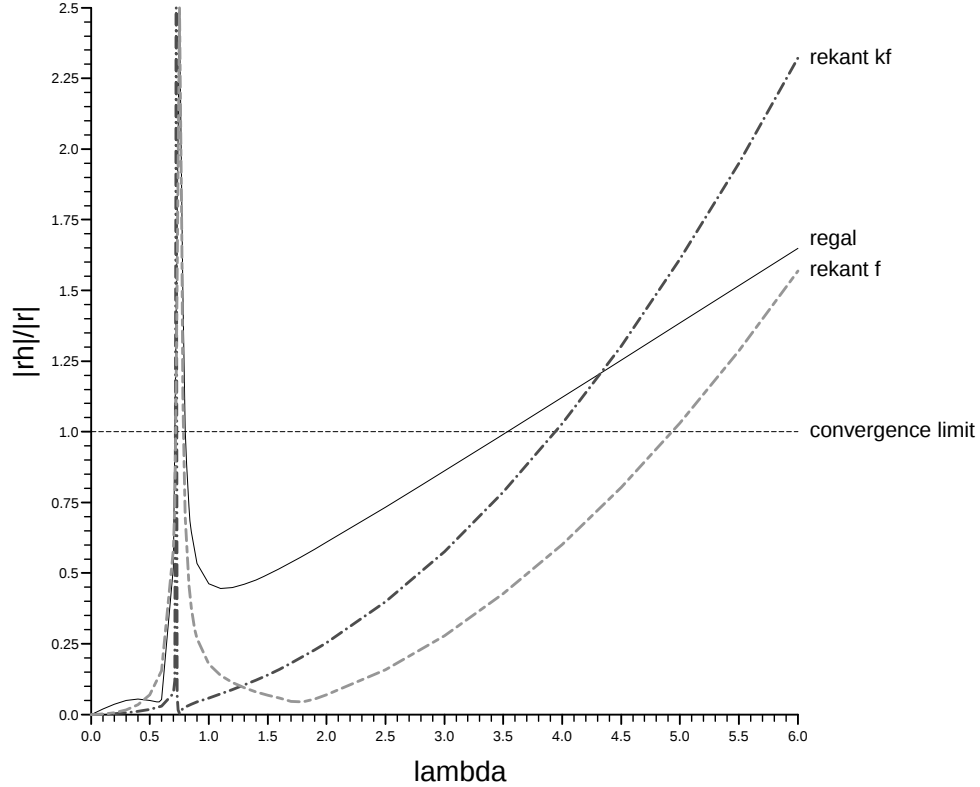


Figure 6.1: Convergence of methods for $k(x, t) = 2\lambda \max(x, t)$

Figure (6.1) illustrates the variation in $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$ (written as $|rh|/|r|$ in the Figure) for each of the methods, with λ taken in the range, $0 \leq \lambda \leq 6$. The most striking feature of the graph is the spike which occurs at around $\lambda = 0.7$. It seems reasonable to suppose that this spike in some way represents the presence of an eigenvalue, since the ratio of the residual norms approximates the spectral radii of the operators S and SK .

There is a process by which it is occasionally possible to calculate, analytically,

Finally,

$$S = (I - KP_n)^{-1}(K - KP_n) = I - (I - KP_n)^{-1}(I - K),$$

so

$$\begin{aligned}(Sf)(x) &= f(x) - (f(x) - (Kf)(x)) + \int_0^1 c\lambda(1+x^2) f(t) - (Kf)(t) dt \\ &= (Kf)(x) + c\lambda(1+x^2) \int_0^1 f(t)dt - \int_0^1 f(t)dt \int_0^1 2\lambda \max(s, t) ds \\ &= \int_0^1 2\lambda \max(x, t) + c\lambda(1+x^2) - c\lambda^2(1+x^2)(1+t^2) f(t)dt,\end{aligned}$$

$$\text{ie. } (Sf)(x) = \int_0^1 2\lambda \max(x, t) + c\lambda(1+x^2)(1-\lambda-\lambda t^2) f(t)dt.$$

To find the eigenvalues of S , let ϕ now satisfy $\phi = \Lambda S\phi$,

$$\begin{aligned}\text{ie. } \phi(x) &= 2\Lambda\lambda \int_0^x x\phi(t)dt + 2\Lambda\lambda \int_x^1 t\phi(t)dt - c\Lambda\lambda(1+x^2) \int_0^1 (\lambda t^2 + \lambda - 1)\phi(t)dt \\ \phi'(x) &= 2\Lambda\lambda\phi(x) - 2c\Lambda\lambda \int_0^1 (\lambda t^2 + \lambda - 1)\phi(t)dt \\ \phi''(x) &= 2\Lambda\lambda\phi(x) - 2c\Lambda\lambda \int_0^1 (\lambda t^2 + \lambda - 1)\phi(t)dt\end{aligned}$$

$$\text{where } \phi'(0) = 0 \quad \phi'(1) = -(1).$$

Now,

$$\begin{aligned}\int_0^1 \phi(x) dx &= \Lambda \int_0^1 \phi(x) dx - \int_0^1 2 \max(x, t) \phi(t) dt - \int_0^1 (\lambda t^2 + \lambda - 1)\phi(t) dt \\ &= \frac{1}{3} \Lambda \int_0^1 (1 + 3t^2) \phi(t) dt\end{aligned}$$

Hence

$$(1 - \frac{1}{3} \Lambda) \int_0^1 \phi(x) dx = \Lambda \int_0^1 t^2 \phi(t) dt$$

and so,

$$2\Lambda \int_0^1 \phi(x) dx = 2 \int_0^1 (1 - \frac{1}{3} \Lambda) t^2 \phi(t) dt - 2 \Lambda \int_0^1 (1 - t^2) \phi(t) dt = 2$$

where $\quad = (1-\Lambda)^{-1}_0(\quad)$.

Suppose $\Lambda > 0$, and write $2\Lambda = \omega^2(\quad \mathbb{R})$.

Then

$$\begin{aligned}(\quad) &= \cos(\quad) + \sin(\quad) \frac{1}{\Lambda} \\ &= \cos(\quad) \frac{1}{\Lambda}\end{aligned}$$

(since $\varphi'(0) = 0$). Therefore

$$\begin{aligned}\frac{1}{0}(\quad) &= -\sin \frac{1}{\Lambda} \\ &= (1-\Lambda)^{\frac{a}{\omega}} \sin \frac{(1-\Lambda)^{\frac{\xi}{\Lambda\lambda}}}{\omega} \\ &\quad \Lambda\text{-Pxj8jCA8ffC}\end{aligned}$$

,

—

—

$$\frac{1-\Lambda}{\omega}$$

$$\frac{1}{2\lambda} \, 2$$

$$\frac{\quad}{\quad} \, 2 \quad \, 2$$

$$\frac{\quad}{2} \quad \frac{\quad}{2}$$

$$n \quad n$$

$$\frac{-1}{n} \quad \frac{1}{2\lambda} \, \frac{2}{n} \quad n \quad \frac{-2\lambda}{\hat{\omega}_n^2} \quad 1$$

$$2$$

$$\frac{\quad}{2} \quad \frac{\quad}{2}$$

Let the roots of this equation be $\tilde{\omega}_n(\tilde{\omega}_n \rightarrow 0)$ for $n \in \mathbb{N}$, and denote the smallest root by $\tilde{\omega}_1$.

The spectral radius of the operator $\mathcal{A}$ is therefore given by

$$\rho(\mathcal{A}) = \frac{2}{\tilde{\omega}_1} \tag{6.4}$$

where $\tilde{\omega}_1 = \min(\hat{\omega}_1, \tilde{\omega}_1)$.

In order to determine the minimum of $\hat{\omega}_1$ and $\tilde{\omega}_1$ it is necessary to explore the behavior of the roots of equations (6.2) and (6.3). Such analysis is quite involved and so only the results are given here.

For $0 \leq \epsilon \leq 0.5$ equation (6.3) has no non-trivial roots, unlike equation (6.2), and so we take $\tilde{\omega}_1 = \hat{\omega}_1$, giving $\rho(\mathcal{A}) = \frac{2\lambda}{\tilde{\omega}_1}$ which increases linearly as ϵ tends to 0.5.

For $0.5 < \epsilon \leq 0.75$ equation (6.3) has one root which starts off being very large, but which tends to zero like $\frac{1}{\epsilon}$ (where $\epsilon \rightarrow 0$) as $\epsilon \rightarrow 0.75$. This results in $\rho(\mathcal{A}) \rightarrow \infty$ as $\tilde{\omega}_1 \rightarrow 0$, giving $\rho(\mathcal{A}) = \frac{2\lambda}{\tilde{\omega}_1}$.

$$\frac{2\lambda}{\tilde{\omega}_1}$$

$$\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$$

eigenvalue of the operator K (ie. for $\lambda = \frac{1}{\mu_n}$, where μ_n is an eigenvalue of K), the numerical method would not be able to find an approximation and hence would not converge. This would result in the ratio of the residual norms, and therefore $\rho(S)$ and $\rho(SK)$, exceeding the convergence limit, and in fact, becoming infinite. If so, the spikes in Figure (6.1) would appear to mark the presence of an eigenvalue of the kernel operator K , given by $\mu = \frac{1}{\lambda}$. Since λ is positive in the examples, and K has only one positive eigenvalue, the graphs seem to provide a means of obtaining the maximum positive eigenvalue of K .

Unfortunately the graphs do not tend to infinity at the same point. The RIG and RIK($\chi = f$) methods both become infinite at $\lambda = 0.75$, whilst the RIK($\chi = Kf$) method becomes infinite at a value closer to $\lambda = 0.72$. The problem is that we are not actually solving $\phi = f + K\phi$, but rather an approximate equation lying in a reduced space. The eigenvalues we are locating, therefore, are not of the original equation, but belong instead to the approximate equation.

The approximate eigenvalues arise from the initial Galerkin method and are unaltered by iteration or re-iteration. Thus, in the one-dimensional case with $\chi = f$, the Galerkin equation is $\alpha\{\|f\|^2 - \lambda(Kf, f)\} = \|f\|^2$, which produces the approximate eigenvalue $\tilde{\mu}_1 = \frac{(Kf, f)}{\|f\|^2}$. Since K is self-adjoint we can deduce that $\mu_1^- \leq \tilde{\mu}_1 \leq \mu_1^+$ by considering the Rayleigh quotient. The eigenvalue apparent in the graphs is therefore an underestimate of the maximum positive eigenvalue of K . When we change the subspace, as for the RIK($\chi = Kf$) method, we also change the value of $\tilde{\mu}_1$, in this case to $\tilde{\mu}_1 = \frac{\|Kf\|^2}{(Kf, f)}$. Since $\frac{(Kf, f)}{\|f\|^2} \leq \frac{\|Kf\|^2}{(Kf, f)}$, this second $\tilde{\mu}_1$ will be no less than (and in practice greater than) the first $\tilde{\mu}_1$. Hence, taking $\chi = Kf$ produces a closer underestimate to the maximum positive

eigenvalue of K than can be achieved by taking $\chi = f$.

This result is echoed in the graphs, where the $\text{RIK}(\chi = f)$ and RIG (also with $\chi = f$) methods offer $\tilde{\mu}_1 = 1 \cdot 33$, corresponding to $\lambda = 0 \cdot 75$, whilst the $\text{RIK}(\chi = Kf)$ method gives $\tilde{\mu}_1 = 1 \cdot 39$, since $1 \cdot 38 \leq \mu_1^+ \leq 1 \cdot 42$ can be obtained by standard eigenvalue estimation methods, as mentioned earlier.

Away from the eigenvalue, both of the RIK methods exhibit better convergence properties than the RIG method, as illustrated by their lower $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$ values, which in turn lead to increased intervals of convergence. Although the RIK graphs intersect the RIG graph, they do not do so until all the methods have exceeded the limit for convergence. For this particular problem we can conclude that the RIK methods exhibit consistently better convergence properties than the RIG method, and that, away from the eigenvalue, taking $\chi = f$ constitutes a better choice of subspace than $\chi = Kf$.

6.2 The kernel $\frac{\lambda}{2\kappa_0} \sin(\kappa_0|x - t|)$

The second equation we consider is the general form of equation (5.2), ie.

$$\phi(x) = \cos(\kappa_0 x) + \lambda \int_0^1 \frac{1}{2\kappa_0} \sin(\kappa_0|x - t|) \phi(t) dt, \quad 0 \leq x \leq 1 \quad (6.5)$$

where κ_0 is a parameter (taken to be $\kappa_0 = 1 \cdot 5$ in the following example).

Unlike equation (6.1), finding both $\rho(S)$ and $\rho(SK)$ analytically proves too difficult, so we have to rely entirely on the numerical results in order to gain information about the convergence properties of the methods for this kernel. Figure (6.2) shows the graphs of $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$ against λ for the three methods when

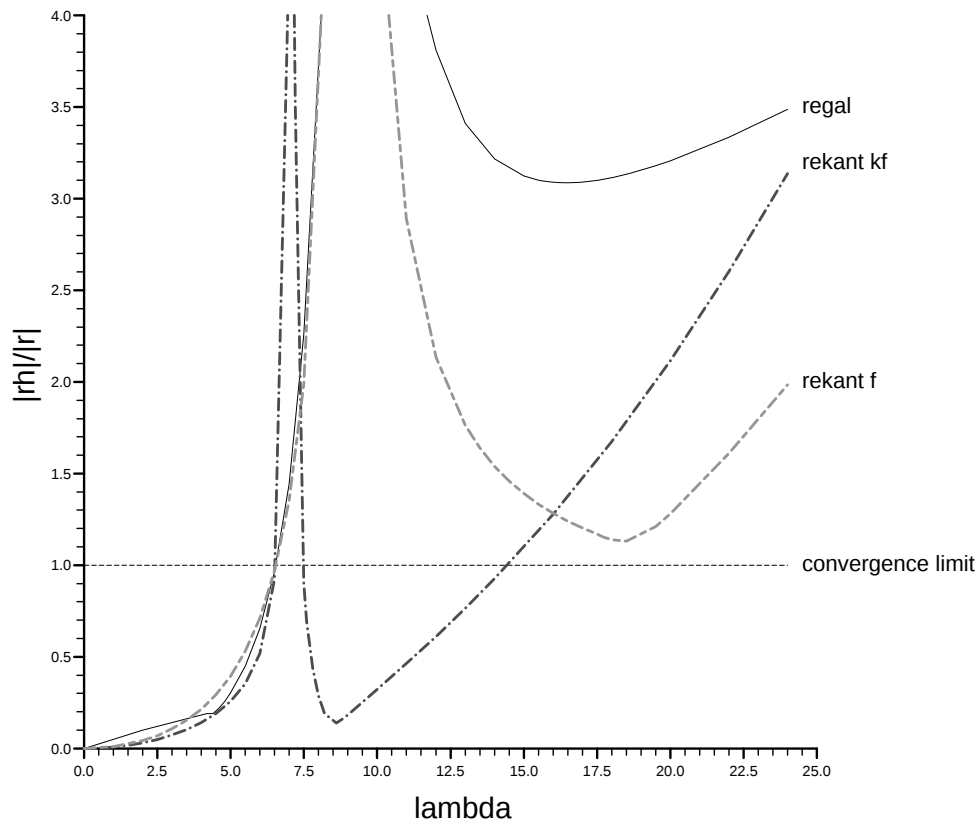


Figure 6.2: Convergence of methods for $(\quad) = \frac{\lambda}{2} \sin(\quad)$

approximating equation (6.5). As before we notice discontinuities in the results, pointing to the presence of approximate eigenvalues. This time, however, the singularities occur at markedly different values for the methods with different subspaces. The $\quad = \quad$ methods become infinite at around $\quad = 9.1$, whilst the $\quad = \quad$ method becomes infinite at approximately $\quad = 7.1$. As in the previous example, we can show that the latter choice of subspace provides the closest underestimate to the maximum positive eigenvalue of $\quad$, giving us $\quad_1 = \frac{1}{\lambda}$
 $0.14 \quad_1^+$.

gives

$$\begin{aligned} \lambda_1^+ &= \min(0.173704 + \frac{\lambda_0^2}{24}, \frac{1}{2} \sqrt{\frac{\lambda_0^2}{2} - \sin^2 \theta_0}) \\ \lambda_1^- &= \max(-0.2, \frac{1}{2} \sqrt{\frac{\lambda_0^2}{2} - \sin^2 \theta_0}) \end{aligned}$$

ie. taking $\theta_0 = 1.5$ we have

$$\begin{aligned} \lambda_1^+ &= 0.17604 \\ \lambda_1^- &= 0.10132 \end{aligned}$$

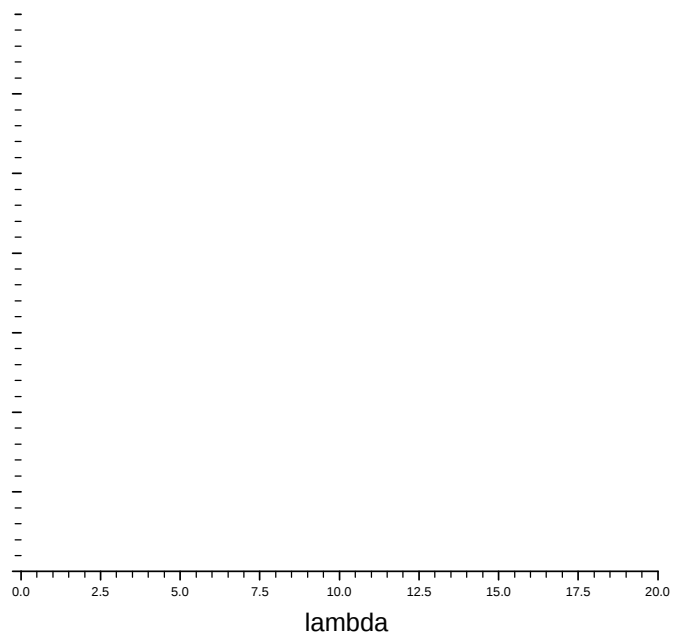
Our approximation, $\tilde{\lambda}_1$, certainly lies below Chamberlain's upper bound, and since it is an underestimate, provides us with a lower bound on λ_1^+ .

All three methods, in this example, demonstrate convergence up to $\theta = 6.5$, where the effect of the approximate eigenvalue becomes too great. Neither the

$$\frac{n}{n-1}$$

$$\frac{n}{n-1}$$

$$n$$



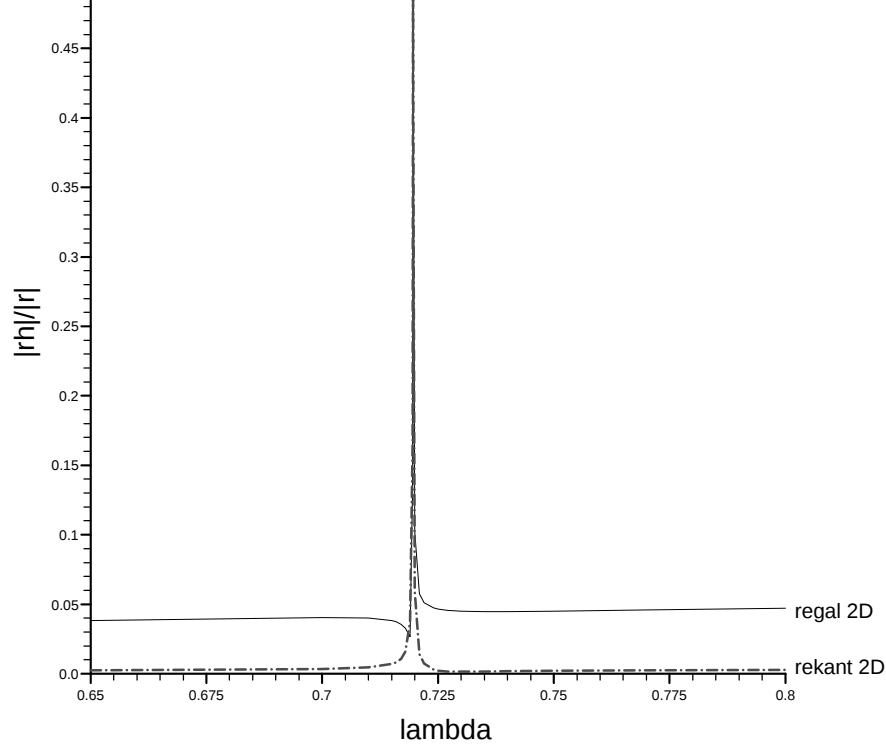
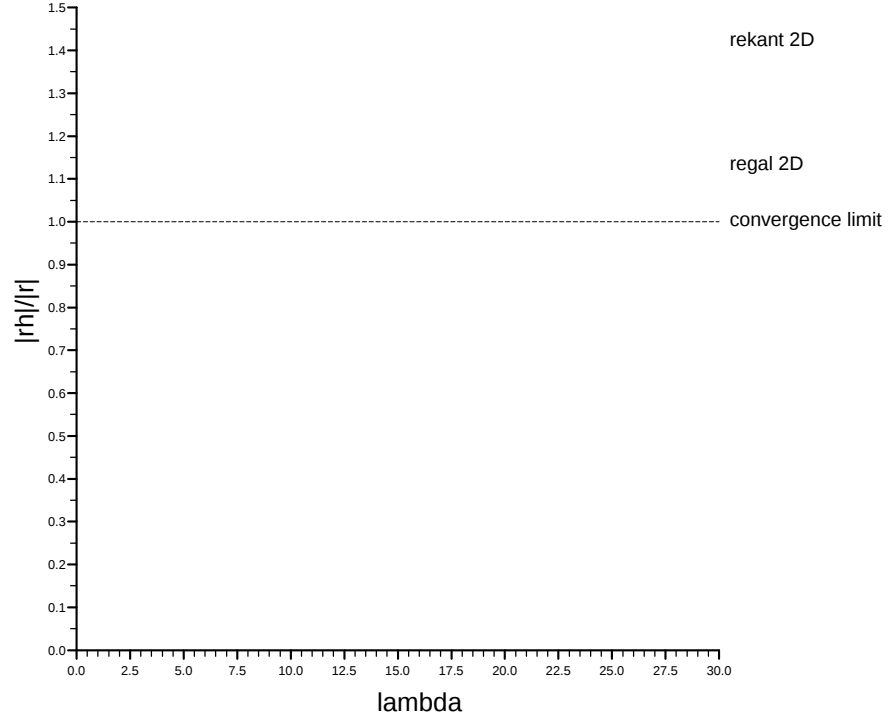


Figure 7.2: The eigenvalue of $(\lambda) = 2 \max(\lambda)$ for a 2-D subspace

magnification in Figure (7.2), the value of λ at which the approximate eigenvalue occurs can be pinpointed at $\lambda = 0.7167$, corresponding to $\lambda_1 = 1.38949$. The disruption caused to the convergence of the methods by the presence of the eigenvector is greatly reduced, allowing us to see more clearly the linear nature of the RIG method, and the quadratic nature of the RIK method. It also demonstrates that the RIG method, which was the most disrupted in the one-dimensional case, may converge faster than the RIK method for certain values of λ , and hence for certain kernels. The point of cross-over in this example does not, however, occur until $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$ is significantly greater than 0.75, a value at which, in a practical situation, we would be looking to utilise an even higher dimensional subspace.

Finally, Figure (7.3) illustrates the effect of moving to a two-dimensional subspace for equation



Chapter 8

Conclusions

The results in the previous two Chapters suggest that the new re-iterated Kantorovich method will converge, and do so faster than the re-iterated Galerkin method, if

$$\rho(SK) < \rho(S) \leq 1.$$

When $\|K\|$ is large, ie. when λ is large in the previous graphs, the RIG method can display the better convergence properties. However, if the view is taken that values of $\frac{\|\hat{r}_n\|}{\|\hat{r}_{n-1}\|}$ greater than 0.5 , corresponding to roughly forty iterations, signal the need for a higher dimensional subspace in which to seek a solution, then the RIK method represents the best practical method of approximation. Naturally, increasing the size of the subspace, and therefore introducing more degrees of freedom into the approximation, has the effect of increasing the rate of convergence, leading to increased regions of convergence.

The presence of approximate eigenvalues can cause a fair amount of disruption to the methods. This is especially true when using a subspace of low dimension, where the number of terms in the expansion of the approximation available to

compensate for the rapid growth of the residuals, is low. It has been shown that choosing a subspace spanned by $\phi_1 = \phi_1^+$, rather than by $\phi_1 = \phi_1^-$, for the one-dimensional case, not only gives a closer underestimate to the exact eigenvalue λ_1^+ of the problem, but also reduces the effect of the approximate eigenvalue on convergence. The position of the approximate eigenvalue may be calculated beforehand and so be avoided. This is done by considering the initial Galerkin equation, given by (in 1-D case)

$$\lambda_1 \int_0^1 \phi_1^2 dx - \left(\int_0^1 \phi_1 dx \right)^2 = \left(\int_0^1 \phi_1 dx \right)$$

hence the eigenvalue, $\tilde{\lambda}_1$, is given by

$$\tilde{\lambda}_1 = \frac{\left(\int_0^1 \phi_1 dx \right)^2}{\int_0^1 \phi_1^2 dx}$$

where ϕ_1 is the choice of trial function. Also to be avoided, naturally, are the

Whereas the estimates of () proved to be very accurate, a similar degree of accuracy could not be attributed to the underestimates of , given by $\frac{\|x_n\|}{\|r_n\|}$ for both methods. These were shown to be very poor in comparison, depending upon the residual, r_n , being a close approximation to a constant multiple of the eigenvector v_1^+ .

There is obviously a need for more exhaustive testing of the re-iterated Kantorovich method, as well as more in-depth analysis, before its strengths and weaknesses are fully understood. In the limited time available, it has only been possible to explore some basic ideas in relation to a few particular examples and many issues have arisen which require closer analysis and examination. Some obvious features that have not been explored in this report include: a closer study

$$\frac{1}{2} \frac{1}{0} \quad 0 \quad 0$$

$$0$$

iterated Galerkin method could be exploited in variational principles. It would

