



An Adaptive Hierarchical Finite Element Method for Modelling Liquid Crystal Devices

Stephen Cornford, Christopher J.P. Newton

HP Laboratories
HPL-2011-143R1

Keyword(s):

Nematic liquid crystals, hierarchic finite elements, adaptive meshes.

Abstract:

Numerical models of liquid crystal devices containing topological defects must take into account two disparate length scales. While the device might extend over several microns, the liquid crystal's alignment varies over nanometres in the region of a defect. Discretising the entire region so finely would be costly in computational terms, but as there are only a few defects, adaptive mesh refinement techniques become attractive. Here, we develop and test an adaptive method which makes use of hierarchical finite elements so that higher order polynomials are used to resolve fine scale features, rather than a finely divided mesh.

External Posting Date: December 8, 2011 [Fulltext]

Approved for External Publication

Internal Posting Date: December 8, 2011 [Fulltext]

An adaptive hierarchical finite element method for modelling liquid crystal devices *

Stephen Cornford^{†‡}

Christopher J.P. Newton[§]

Abstract

Numerical models of liquid crystal devices containing topological defects must take into account two disparate length scales. While the device might extend over several microns, the liquid crystal's alignment varies over nanometres in the region of a defect. Discretising the entire region so finely would be costly in computational terms, but as there are only a few defects, adaptive mesh refinement techniques become attractive. Here, we develop and test an adaptive method which makes use of hierarchical finite elements so that higher order polynomials are used to resolve fine scale features, rather than a finely divided mesh.

Keywords

Nematic liquid crystals, hierarchic finite elements, adaptive meshes.

1 Introduction

Nematic liquid crystals are made up of elongated molecules which tend to align with one another, but are otherwise unordered. Provided that the alignment is uniaxial, and the degree of order does not change, it is usual to denote the average alignment of these molecules over a small volume by a unit vector $\mathbf{n}$, known as the director. In a conventional display device, a layer of liquid crystal is sandwiched between transparent plates whose surfaces are treated to influence the director adjacent to them. Elastic forces then cause the liquid crystal to adopt a ground state in which the director varies smoothly from one surface to the other. Applying an electric field changes the orientation of the director, allowing the optical properties of the cell to be controlled. In this kind of device the director returns to its ground state once the electric field has been removed.

Among recent innovations in liquid crystal display technology are bistable nematic devices, in which both a light and a dark state are stable in the absence of an electric field. In at least three of these devices: the Zenithal Bistable Device (ZBD) [1], the Post-Aligned Bistable Nematic device (PABN) [2], and a device based upon rectangular wells [3], only the presence of discontinuities in $\mathbf{n}$ - line defects - permits the device to be bistable at all. More than that, switching these devices

*This work was supported in part by EPSRC award EP/E040993/1.

[†]Department of Mathematics and Statistics, University of Strathclyde, Glasgow G1 1XH, UK.

[‡]Now at The School of Geographical Sciences, University of Bristol, University Road, Bristol BS8 1SS, UK. eltadcirad@gmail.com.

[§]Hewlett-Packard Laboratories, Long Down Avenue, Stoke Gifford, Bristol BS34 8QZ, UK. chris.newton@hp.com.

between their states involves the dynamics of line defects, so it is essential to be able to include them in calculations.

Both the bulk of the sample of nematic liquid crystal and any defects in it can be described without introducing singularities by a symmetric, traceless second rank alignment tensor, $\mathbf{Q}$, often called the $\mathbf{Q}$ -tensor. Numerical models based upon the alignment tensor have been studied for two decades. Early work concentrated on the study of defects [4, 5] while realistic bistable devices have been studied in recent publications [8, 9, 10]. Three kinds of alignment are accounted for, depending on the eigenvalues of $\mathbf{Q}$. Wherever the eigenvalues are all equal (and so zero), the liquid crystal is isotropic: there is no order, and no direction of alignment. If only two of the eigenvalues are equal the liquid crystal is uniaxial, and the eigenvector associated with the distinct eigenvector can be equated with the director. If all three eigenvalues are different, the liquid crystal is biaxial, with different order along two orthogonal directions. In a realistic device, the liquid crystal is *positive uniaxial*, with one positive eigenvalue larger in magnitude than the other two, nearly everywhere, but in the neighbourhood of a defect a biaxial region encloses a negative uniaxial central point.

The size of this biaxial and negative uniaxial region is determined by competition between two energy densities, a thermotropic, or Landau-de Gennes energy density

$$\omega_B = a \text{tr}(\mathbf{Q}^2) + \frac{2b}{3} \text{tr}(\mathbf{Q}^3) + \frac{c}{2} (\text{tr}(\mathbf{Q}^2))^2, \quad (1)$$

and an elastic energy density, which in its simplest form is

$$\omega_F = \frac{L_1}{2} \sum_{ijk} \left(\frac{\partial Q_{ij}}{\partial x_k} \right)^2. \quad (2)$$

Of these, the thermotropic energy depends only on the eigenvalues of $\mathbf{Q}$, λ_1, λ_2 and λ_3 say. The coefficients a, b, c are such that ω_B has its minimum when the liquid crystal is positive uniaxial, with an order parameter

$$\begin{aligned} S &= \frac{3}{2} \lambda_1 = -3\lambda_2 = -3\lambda_3 \\ &= \frac{1}{4c} \left[-b + (b^2 - 24ac)^{1/2} \right]. \end{aligned} \quad (3)$$

The parameter L_1 is related to the more usual one-constant elastic modulus k by $L_1 = k/(2S_e^2)$, where S_e is the value of S at which k was measured. Here, we use values for 5CB at 4K below the pseudo-critical temperature T^* (where the isotropic state loses stability): $a = -0.39 \text{ MJm}^{-3}$, $b = -3.6 \text{ MJm}^{-3}$, $c = 4.4 \text{ MJm}^{-3}$ [13]¹, $S_e = 0.62$ and $k = 6.0 \text{ pN}$. Looking at these, we see that ω_B and ω_F only become comparable, so that the liquid crystal is forced away from the phase given by (3), when $\mathbf{Q}$ varies over length scales of $\sim 10 \text{ nm}$.

Compared to the extent of a typical device - a few microns - 10 nm is a tiny length scale, so we are motivated to investigate adaptive methods. If we were to discretise the whole of the computational domain finely enough to resolve defects, the number of degrees of freedom in the resulting problem would be huge. However the number of defects is usually small, so one could

¹The definitions of $\mathbf{Q}$ and the coefficients in (1) differ, by simple factors, between this paper and [13]. Compared to A, B , and C given there, $a = 3/4A(T - T^*)$, $b = 9/4B$ and $c = 9/8C$.

perhaps discretise the domain finely enough around them provided that they can be detected in the first place on a coarse mesh. There is a complication, though: in simulations where the defect moves, it will be necessary to generate many meshes as the defects move, and somehow transfer the solution between them. To this end, we investigate the use of hierarchical bases [14, 15], where, rather than resolving defects by generating a mesh with many small elements close to them, we introduce higher order polynomials into large elements.

2 Method

As the focus of this paper is the development of an adaptive method which copes with the disparate length scales discussed, we consider a rather simple, two-dimensional steady state model. A possibly periodic domain Ω is bounded by a contour Γ , and over this domain the free energy

$$F = \int_{\Omega} (\omega_F + \omega_B + \omega_E) \, d\Omega + \int_{\Gamma} \omega_S \, d\Gamma \quad (4)$$

is minimised. The first integral in (4) is comprised of terms due to the thermotropic and elastic energies, and a dielectric term which reflects the tendency of the liquid crystal to be aligned by an electric field $\mathbf{E}$

$$\omega_E = -\epsilon_0 \mathbf{E} \cdot \left(\frac{\epsilon_a}{S_e} \mathbf{Q} + \frac{\epsilon_a + 3\epsilon_{\perp}}{3} \mathbf{I} \right) \cdot \mathbf{E}. \quad (5)$$

Here $\mathbf{I}$ is the 3×3 identity matrix, ϵ_0 is the permittivity of free space, $\epsilon_{\perp}$ is the relative dielectric permittivity perpendicular to the director, while ϵ_a is the relative dielectric anisotropy. We use the values $\epsilon_{\perp} = 7$ and $\epsilon_a = 11$ for the relative permittivities in (5). The second integral is a result of the forces which tend to align the liquid crystal at a solid surface. We set the surface energy density to

$$\omega_S = \frac{W}{2} \text{tr}((\mathbf{Q}_S - \mathbf{Q})^2). \quad (6)$$

where $\mathbf{Q}_S$ is the value which $\mathbf{Q}$ would adopt at the surface in the absence of other forces, and W is the anchoring strength.

In this initial investigation, we have omitted some potentially important contributions to the free energy, such as those due to the flexoelectric effect which can account for switching between states in bistable devices [7, 8], because it is the coupling between the elastic and thermotropic energies which leads to the need for adaptive techniques. For the same reason, we will just consider problems with a known uniform electric field.

Since $\mathbf{Q}$ is symmetric and traceless, it can be parameterised by only five components. We make use of the formulation

$$\mathbf{Q} = \frac{1}{\sqrt{2}} \begin{pmatrix} q_1 - \frac{1}{\sqrt{3}}q_0 & q_2 & q_3 \\ q_2 & -q_1 - \frac{1}{\sqrt{3}}q_0 & q_4 \\ q_3 & q_4 & \frac{2}{\sqrt{3}}q_0 \end{pmatrix} \quad (7)$$

as used in [4, 5]. So, to minimise (4), we must solve to five Euler-Lagrange equations in q_i ,

$$\nabla \cdot (L_1 \nabla q_i) - \frac{\partial \omega_B}{\partial q_i} - \frac{\partial \omega_E}{\partial q_i} = 0 \quad \text{on } \Omega, \quad (8)$$

together with boundary conditions chosen from one of two types: mixed boundary conditions derived from the surface energy, which give

$$L_1 \boldsymbol{\nu} \cdot \boldsymbol{\nabla} q_i = W(q_{Si} - q_i) \quad \text{on } \Gamma, \quad (9)$$

or periodic boundary conditions.

2.1 Finite element formulation of the Q -tensor model

The equations (8) are nonlinear because of the thermotropic term. We could follow one of two equivalent routes to solve them numerically: either linearise the nonlinear PDEs (8) to find a sequence of linear PDEs, and discretise these, or discretise the nonlinear equations first and then linearise those. Following the first route, we set $q_i = \bar{q}_i + \delta q_i$, and discard higher order terms in δq_i , which leads to linear equations in δq_i

$$\boldsymbol{\nabla} \cdot (L_1 \boldsymbol{\nabla} \delta q_i) - \sum_j \omega_{B,ij}(\bar{\mathbf{q}}) \delta q_j = -\boldsymbol{\nabla} \cdot (L_1 \boldsymbol{\nabla} \bar{q}_i) + \omega_{B,i}(\bar{\mathbf{q}}) + \left. \frac{\partial \omega_E}{\partial q_i} \right|_{\bar{q}_i} \quad \text{on } \Omega \quad (10)$$

and

$$\boldsymbol{\nu} \cdot (L_1 \boldsymbol{\nabla}(\delta q_i)) + W \delta q_i = -\boldsymbol{\nu} \cdot (L_1 \boldsymbol{\nabla}(\bar{q}_i)) + W(q_{Si} - \bar{q}_i) \quad \text{on } \Gamma \quad (11)$$

where

$$\omega_{B,i}(\bar{\mathbf{q}}) = \left. \frac{\partial \omega_B}{\partial q_i} \right|_{\bar{\mathbf{q}}} \quad (12)$$

and

$$\omega_{B,ij}(\bar{\mathbf{q}}) = \left. \frac{\partial^2 \omega_B}{\partial q_i \partial q_j} \right|_{\bar{\mathbf{q}}}. \quad (13)$$

Repeatedly solving these for δq_i and then updating $\bar{q}_i$ until the right hand sides vanish is, of course, Newton's method.

Using Galerkin's method to discretise the linear PDEs (10) and boundary conditions (11) is a well worn path [16], so we only note the one point which affects our discussion: we can choose any of a host of discretisation schemes (with different meshes, for example) by choosing a function space X with N basis functions ϕ_μ , and writing components of the solution as

$$\bar{q}_i = \sum_{\mu=1}^N \bar{u}_{(\mu+iN)} \phi_\mu \quad (14)$$

and

$$\delta q_i = \sum_{\mu=1}^N \delta u_{(\mu+iN)} \phi_\mu. \quad (15)$$

The vectors $\bar{\mathbf{u}}$ and $\delta \mathbf{u}$ have $5N$ components, or *degrees of freedom*. In the end, we have a sparse linear system

$$[K + A(\bar{\mathbf{u}}) + S] \delta \mathbf{u} = -[K \bar{\mathbf{u}} + a(\bar{\mathbf{u}}) + d + s] \quad (16)$$

or, gathering terms together,

$$J(\bar{\mathbf{u}})\delta\mathbf{u} = -f(\bar{\mathbf{u}}) \quad (17)$$

which must be solved for the vector $\delta\mathbf{u}$ at each iteration of Newton's method. Currently we are using the PETSc SNES nonlinear solver [17] alongside the direct linear system solver MUMPS [18] to perform the Newton iterations. To aid convergence we have introduced a continuation parameter σ multiplying the thermotropic terms. This parameter increases from a value $\sigma_0 < 1$ to 1 as the iterations progress. This does not alter the final solution, but widens the radius of convergence around it, see Section 2.6.

The left hand side of (16) is made up from a stiffness matrix with elements

$$K_{(\mu+iN)(\nu+iN)} = \int_{\Omega} L_1 \nabla \phi_{\mu} \cdot \nabla \phi_{\nu} d\Omega, \quad (18)$$

a mass-like matrix of thermotropic contributions with elements

$$A_{(\mu+iN)(\nu+jN)}(\bar{\mathbf{u}}) = \sigma \int_{\Omega} \omega_{B,ij}(\bar{\mathbf{q}}) \phi_{\mu} \phi_{\nu} d\Omega, \quad (19)$$

and a matrix due to the surface conditions

$$S_{(\mu+iN)(\nu+iN)} = W \int_{\Gamma} \phi_{\mu} \phi_{\nu} d\Gamma. \quad (20)$$

The residual vector on the right hand side of (16) is the sum of an elastic term $K\bar{\mathbf{u}}$, a thermotropic term, with components

$$a_{\mu+iN}(\bar{\mathbf{u}}) = \sigma \int_{\Omega} \phi_{\mu} \omega_{B,i}(\bar{\mathbf{q}}) d\Omega, \quad (21)$$

a dielectric term

$$d_{\mu+iN} = \int_{\Omega} \phi_{\mu} \frac{\partial \omega_E}{\partial q_i} d\Omega, \quad (22)$$

and a surface term

$$s_{\mu+iN} = W \int_{\Gamma} \phi_{\mu} (q_{Si} - \bar{q}_i) d\Gamma. \quad (23)$$

2.2 Piecewise linear elements

We indicated earlier that the final step in discretisation of the governing equations is the selection of basis functions, which, in finite element methods, means the choice of a particular type of element. For $\mathbf{Q}$ -tensor problems, each of the components q_i and their derivatives must be square-integrable everywhere in the domain if the integrals in (18) are to be well-defined. In other words, the q_i , and hence the bases ϕ_{μ} , must belong to the function space $H^1(\Omega)$. A common choice (in two-dimensional problems) is the triangular piecewise linear element, or first order Lagrange element. The domain Ω is subdivided into triangular regions and a basis functions is defined for each vertex of the resulting mesh. So, if there are N vertices on the mesh then there will be N basis functions. Each basis function is associated with a vertex on the mesh and is obtained by adding contributions

from each triangle containing that vertex. For a given triangle we can label the vertices 1, 2 and 3 and associate with each of these vertices a barycentric coordinate $\lambda_i, i = 1 \dots 3$, where λ_i is 1 at vertex i , 0 at the other two vertices and varies linearly between them. The vertex function, ϕ_μ , is then

$$\phi_\mu = \begin{cases} \lambda_i & \text{if } \mu \text{ is vertex } i \text{ of the triangle,} \\ 0 & \text{otherwise.} \end{cases} \quad (24)$$

Several authors have reported on the use of this kind of element [4] to solve $\mathbf{Q}$ -tensor problems. Others have made use of second order Lagrange elements, where six quadratic functions in the λ_i replace (24).

Once the basis functions ϕ_μ have been chosen, the integrals (18)-(23) can be evaluated. Since a given ϕ_μ is only non-zero inside those triangles with the appropriate vertex, only a few pairs (ϕ_μ, ϕ_ν) contribute non-zero terms to the matrices K, A , and S (which are therefore sparse), and the evaluation of these can be reduced to integration over only a few triangles. These lesser integrations can be approximated by quadrature formulae, such as the common Gauss-Legendre formulae, or, for higher order elements, Grundmann-Möller formulae [11].

If a solution computed on a mesh of piecewise linear elements is not accurate enough, there are four basic options:

- Uniform h -refinement. The very different length scales in this problem make this option impractical, or at least very inefficient.
- Uniform p -refinement. Here every element in the mesh is replaced by one with quadratic, cubic, or higher order and the solution is recalculated. Again the different length scales in the problem make this impractical and inefficient.
- Non-uniform h -refinement. This can be employed, by generating a new mesh where the density of elements varies and is chosen to best resolve the solution [12]. Non-uniform h -refinement seems well suited to the $\mathbf{Q}$ -tensor problem, provided a mesh concentrated close to defects can be generated, but there is a complication. If the defect moves, a mesh, or sequence of meshes, fine enough to capture its entire path must be used. Single meshes have been employed in the study of the ZBD, where the defects remain close to the surface so that the number of refined elements is limited, but could be expensive in cases where defects could move freely in the bulk. Using a sequence of meshes, refining and coarsening parts of the mesh as the defects move is a possibility, but made difficult by the need to interpolate the $\mathbf{Q}$ -tensor profile from one mesh to the next.
- Non-uniform p -refinement. This is also well suited to the problem as refinement can be restricted to elements close to any defect regions. With an appropriate choice of basis functions elements can be easily refined, or coarsened, as necessary to track any defect regions. Interpolation onto the new mesh is also straightforward.

Some combination of these options could also be used, but in the work presented here we have focussed on non-uniform p -refinement using hierarchical basis functions.

2.3 Hierarchical bases and p -refinement

One major difference between the finite element methods reported in this paper and those used elsewhere to solve $\mathbf{Q}$ -tensor problems is our choice of hierarchical bases for the function spaces X . In this context, hierarchical means that for any two of the spaces X and Y that we consider, it is easy to construct a third space Z which satisfies $X \subset Z$ and $Y \subset Z$. This turns out to be a useful property for two reasons: firstly, it is relatively simple to refine or coarsen the solution non-uniformly, and secondly, it leads to an effective method of error estimation.

It is possible to construct H^1 conforming hierarchical bases starting from the first order triangular elements described earlier. Higher order polynomials in the barycentric coordinates λ_μ are progressively added to the basis, so that at a given order p the bases span the space of polynomials of order p in the λ_μ . The particular progression we use has been described by Ainsworth and Coyle, details can be found in [15, 14]. This set of basis functions is designed so that each basis function is associated with either a vertex, an edge, or the interior of an element of the mesh. For every order $p > 1$, an additional basis function of order p is associated with each of the edges of the mesh. Except at the boundary, each of these basis functions gets contributions from the two triangles which share this edge. Let the $\phi_p^{\gamma_k}$ be the order p basis function associated with the edge γ_k , then

$$\phi_p^{\gamma_k} = \begin{cases} \frac{-4}{p(p-1)} \lambda_i \lambda_j L'_{p-1}(\lambda_j - \lambda_i) & \text{if } i \text{ and } j \text{ are vertices of edge } \gamma_k, \\ 0 & \text{otherwise.} \end{cases} \quad (25)$$

where L_{p-1} is the Legendre polynomial of order $p-1$. For every order $p > 2$ additional basis functions of order p are associated with the interior of each triangle. These functions are associated with just one triangle and are zero elsewhere. Let $\phi_{0,0}^e$ be the first of these functions with order 3 and associated with triangle e then

$$\phi_{0,0}^e = \beta = \begin{cases} \lambda_1 \lambda_2 \lambda_3 & \text{for } \lambda_1, \lambda_2, \lambda_3 \text{ defined on } e, \\ 0 & \text{otherwise.} \end{cases} \quad (26)$$

This is then followed by additional functions at fourth order and above. Each increase in p results in adding extra functions $\phi_{i,j}^e$ given by

$$\phi_{i,j}^e = \beta \lambda_2^i \lambda_3^j \quad i + j = p - 3. \quad (27)$$

In all, an order p element contributes to 3 vertex functions, $3(p-1)$ edge functions and $(p-1)(p-2)/3$ interior functions.

There is an obvious drawback to the use of hierarchical bases. Since the number of basis functions associated with each element increases like p^2 , the number of non-zero contributions to the integrals (18)-(23) increases like p^4 . This is to be compared with a conventional basis, where the number of non-zero contributions only increases linearly with the number of elements. So there will be both a cost in CPU time caused by the growing number of evaluations, and more memory will be needed to store the denser linear system. Furthermore, the order of the integrands increase, so that they need to be evaluated at more points for the integral to be approximated with sufficient accuracy. In particular, the matrix of thermotropic contributions (19) has order $4p$, although we have found in practice that a quadrature formula of degree $3p+1$ suffices.

2.4 Refining and coarsening the mesh

It is simple to perform a non-uniform p -adaption of a mesh constructed from the Ainsworth-Coyle hierarchical elements, refining some parts of the mesh and coarsening others. There are two requirements to satisfy: the new space must belong to $H^1(\Omega)$, and it must be possible to transfer the solution from the original space to the new space. Satisfying the first requirement is simple: whenever two elements of order p and $p' > p$ share an edge, the coefficients of every edge function of order greater than p on that edge must be set to zero.

Our second requirement is important only because we are solving a nonlinear problem (but will also be important when we come to solve time-dependent problems). We want to transfer the solution $\mathbf{u}_X$ defined on the original mesh with $\phi_\mu \in X$ to the new mesh with $\phi_\mu \in Y$. First, we find the space $Z = X \cup Y$, and as $X \subset Z$ simply complement the vector $\mathbf{u}_X$ with zeros for each $\phi_\mu \notin X$. Then we form a vector $\mathbf{u}_{XY}$ from part of this vector, retaining only coefficients of $\phi_\mu \in Y$. Note that $\mathbf{u}_{XY}$ is not the solution to the finite element problem defined on the new mesh, but can be used to start a new sequence of Newton iterations, culminating in a refined solution $\mathbf{u}_Y$.

2.5 Error estimation

Whatever method is used to refine the solution, we need to decide where to do so, and it is natural for this decision to be based on estimates of the error in the solution. Having chosen to make use of hierarchical spaces, a type of *a posteriori* error estimation becomes available to us [6]. The idea is to compute a solution $\mathbf{u}_X$ in a function space X , and then estimate the solution $\mathbf{u}_Y$ in an enhanced space, Y , where $X \subset Y$. That done, the difference between these two, $\mathbf{e} = \mathbf{u}_Y - \mathbf{u}_X$, is an estimate of the error in $\mathbf{u}_X$.

The great disadvantage of this kind of error estimation is its potential cost: solving the problem in the larger space Y will be in general more expensive than solving the original problem. To alleviate this, we do not solve the full nonlinear problem in the larger space, but instead compute a solution to the linear system

$$J_Y \delta \mathbf{u}_Y = -\mathbf{f}_X(\bar{\mathbf{u}}_{XY}). \quad (28)$$

Here, the subscript Y denotes assembly of the linear system in the space Y . In other words, we compute the first Newton step $\delta \mathbf{u}_Y$ for the larger problem, starting from the solution of the original problem, and use this as our error estimate. But, although this reduces the computational time needed to compute the error estimate, it does not reduce the extra storage required by the matrix J_Y .

2.6 Automatic p -adaption

Once we have the error estimate $\delta \mathbf{u}_Y$, we can make use of it to decide which elements to refine or coarsen and by how much. A natural measure of the difference between two solutions is the difference in their bulk free energies. Writing $\omega_T = \omega_F + \omega_B + \omega_E$, we define a function

$$F(\mathbf{e}, \mathbf{u}) = \int_e \omega_T(\mathbf{u}) \, d\Omega \quad (29)$$

where the integration limit e denotes the whole of one element. We then calculate a relative energy for each element e ,

$$r(e) = \frac{|F(e, \mathbf{u}_X) - F(e, \mathbf{u}_{XY} + \delta \mathbf{u}_Y)|}{\sum_e |F(e, \mathbf{u}_X)|}, \quad (30)$$

and ideally would like to define new function spaces X and Y so that, in every element, $r(e)$ is less than some tolerance ϵ_r . However, we only know that $r(e)$ should decrease as the order of elements increases, so we implement an iterative procedure. Figure 1 shows the algorithm that we use. The continuation parameter σ multiplies the thermotropic terms. It increases from an initial value $\sigma_0 < 1$ to 1 as the iterations progress. This does not alter the final solution, but widens the radius of convergence around it.

```

Choose positive integers  $n$  and  $m < n$  and positive reals  $\sigma_0$ ,  $\epsilon_n$  and  $\epsilon_r$ , a function space  $X$ 
whose elements have order  $p_X(e)$  and initial solution  $\bar{\mathbf{u}}_X$ 
 $\sigma \leftarrow \sigma_0$ 
for  $i = 0$  to  $n$  do                                     ▷ Outer iteration
  while  $\|\mathbf{f}(\bar{\mathbf{u}})\| > \epsilon_n$  do                             ▷ Newton iteration
    solve  $J_X \delta \mathbf{u}_X = -\mathbf{f}_X(\bar{\mathbf{u}}_X)$ 
     $\bar{\mathbf{u}}_X \leftarrow \bar{\mathbf{u}}_X + \delta \mathbf{u}_X$ 
  end while
   $\sigma \leftarrow \min(1, \sigma_0 + (i + 1)(1 - \sigma_0)/m)$ 
  define  $Y$  such that  $p_Y(e) = p_X(e) + 1$ 
  solve  $J_Y \delta \mathbf{u}_Y = -\mathbf{f}_Y(\bar{\mathbf{u}}_{XY})$ 
  find  $q(e)$  such that  $2^{q(e)} < r(e)/\epsilon_r < 2^{q(e)+1}$ 
  define  $Z$  by  $p_Z(e) = p_X(e) + q(e)$ 
   $X \leftarrow Z$ 
   $\bar{\mathbf{u}}_X \leftarrow \bar{\mathbf{u}}_{XZ}$ 
end for

```

Figure 1: The p -adaption algorithm

3 Results and discussion

As there are two length scales our problems, it makes sense to look at them individually before considering a realistic problem. To that end, we report on the performance of uniform p - and h - refinement for two simple test problems. In the first, we find that p -refinement unambiguously outperforms h -refinement in resolving micron scale variation of the eigenvectors of $\mathbf{Q}$. In the second, we see that while p -refinement can be used to resolve defects, attention must be paid to the mesh spacing as well. We then apply our automatic p -adaption algorithm to a problem with both length scales.

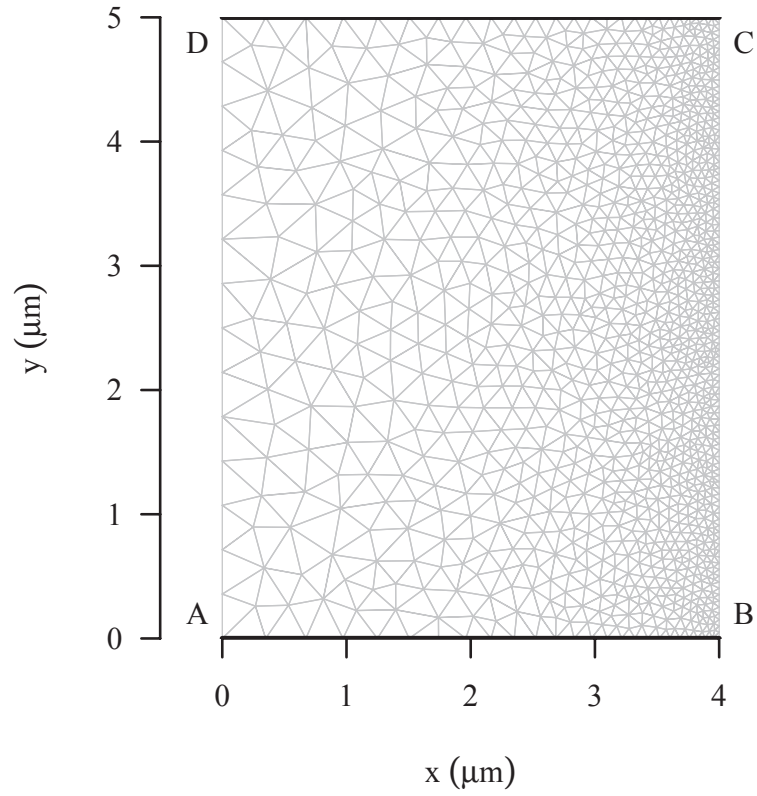


Figure 2: Graded mesh on a rectangle ABCD. Homeotropic alignment is imposed along AB, planar alignment along CD, and free ($W = 0$) boundary conditions along BC and DA. The mesh is smoothly refined from 300 nm on the left hand side to 50 nm on the right.

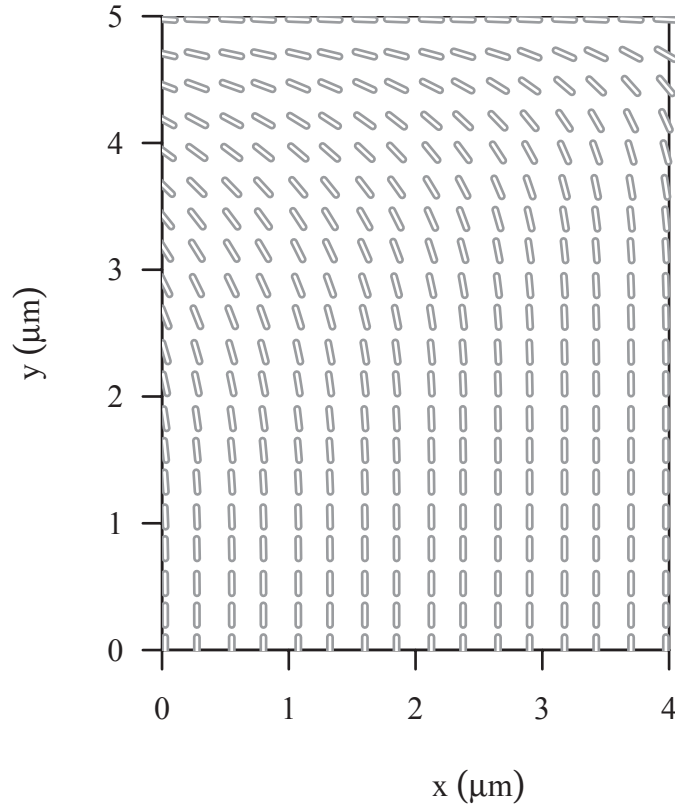


Figure 3: Solution for the HAN problem using $p = 1$ elements. The line segments are parallel to the director. On the left hand side, where the mesh is coarsest, $\mathbf{Q}$ varies in y across most of the cell, while on the right side, where the mesh is fine, $\mathbf{Q}$ only varies significantly in the upper half of the domain.

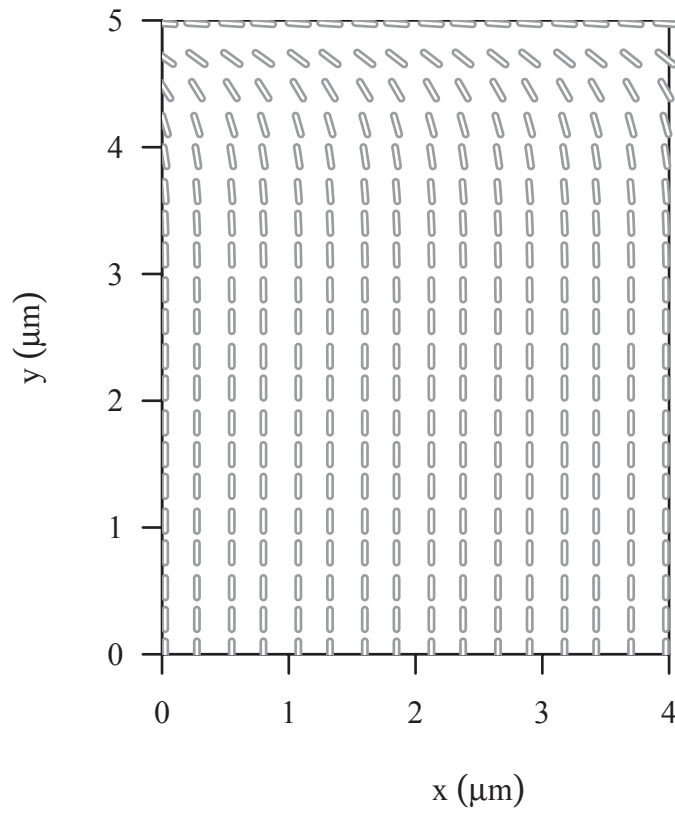


Figure 4: Solution for the HAN problem using $p = 2$ elements. In this case, the solution does not vary noticeably as the mesh is refined from the left to the right, with most of the variation in Q confined to the upper half of the domain

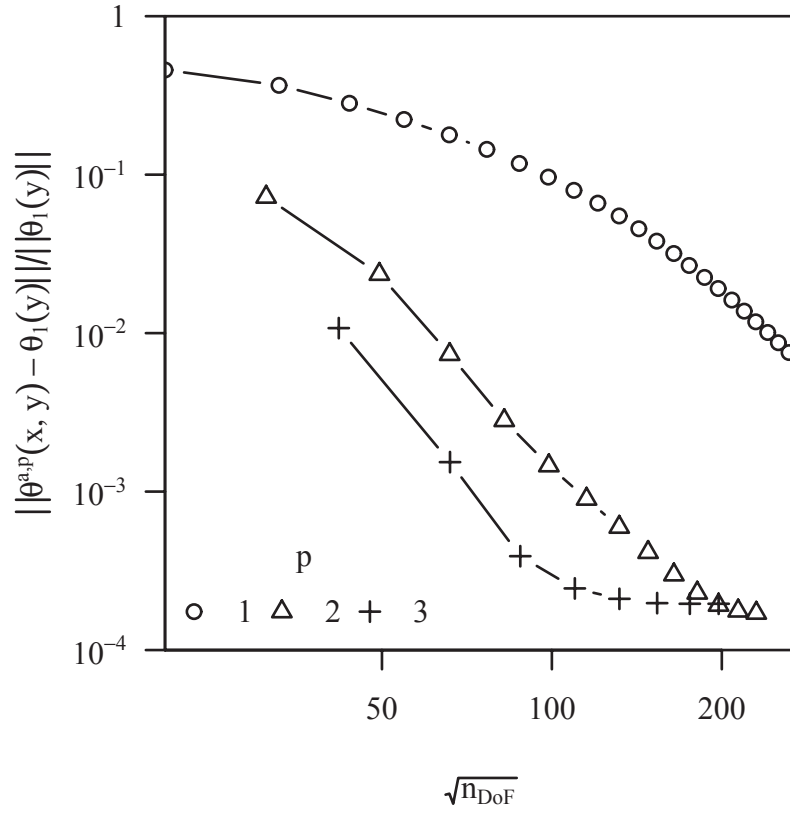


Figure 5: Error in the tilt angle $\theta(x, y)$ plotted against number of degrees of freedom. Following any of the lines from left to right, the mesh spacing $1/a$ is reduced while the polynomial order p is kept constant. For a given number of degrees of freedom (n_{DoF}), the higher p (and so lower h) solution is more accurate.

3.1 Uniform h - and p -refinement without defects

In our first test problem we consider a hybrid-aligned nematic (HAN) cell. These consist of a several microns thick layer of nematic liquid crystal sandwiched between two plane surfaces. One surface is treated to promote homeotropic alignment, where the director prefers to lie perpendicular to the surface. The other surface is treated to promote planar homogeneous alignment, where the director prefers to lie in a defined direction parallel to the surface (in this test problem this preferred direction is the x -axis). Because the cell's surfaces are much larger than its vertical size, it could be treated as a one-dimensional system, although in this case we shall treat it as a two-dimensional problem.

HAN cells are often chosen as experimental and theoretical models of the more complex bistable systems [7], and as such, their behaviour is well known. In particular, they can be described by a simpler model, which only involves $\theta_1(y)$, the angle made between the director and the x -axis, over the vertical extent of the cell $0 \leq y \leq d$. Here we assume that the electric field is constant across the cell

$$\mathbf{E} = (0, \frac{V}{d}, 0) \quad (31)$$

where V is the applied voltage. One then has

$$k \frac{d^2 \theta_1}{dy^2} + \frac{1}{2} \epsilon_0 \epsilon_a \left(\frac{V}{d} \right)^2 \sin 2\theta = 0 \quad (32)$$

with $\theta_1(0) = \pi/2$ and $\theta_1(d) = 0$.

When $V = 0$, (32) has a very straightforward solution

$$\theta_1(y) = \frac{\pi(d-y)}{2d}, \quad (33)$$

so that the director rotates uniformly from the bottom of the cell to the top.

When $V \neq 0$, the director prefers to be aligned parallel to the field, so that as V increases more and more of the cell is vertical. For the results presented here $V = 3$ volts.

If we choose to discretise the $\mathbf{Q}$ -tensor problem with $p = 1$ elements, we must use a rather fine mesh. This can be illustrated by choosing a mesh whose spacing varies in x , such as that shown in Figure 2. Everything about this problem - the boundary conditions, the electric field, the geometry - should result in a solution which varies only in y . But the solution, shown in Figure 3, exhibits a marked dependency on x . Close to the edge AD, where the mesh is coarsest, the solution varies evenly in y across much of the cell, whereas on the opposite edge the solution is much more vertical, as it should be. On the other hand, if we choose $p = 2$ elements - see Figure 4 - then there is little apparent dependence on mesh spacing. Figure 4 shows the solution in this case and it does not vary noticeably in x . All of the variation in y is confined to the upper part of the cell (much as on the fine mesh part of the $p = 1$ solution).

We can look at the relative efficiencies of h - and p -refinement for this HAN problem by generating a sequence of uniform meshes and comparing the solutions found on them with that of the simplified model (32). For these tests the computational domain, Ω is bounded by $0 < x < 1$ and $0 < y < 5$ (measured in microns). Planar homogeneous anchoring is imposed at $y = 0$ and homeotropic anchoring at $y = 5$, while periodic conditions are applied at $x = 0$ and $x = 1$. Each

mesh was generated by subdividing Ω into squares $1/a$ microns on a side, and subdividing each of those into two right angled triangles. Solving the finite element problem on a given mesh, distinguished by a , and for a given polynomial order p , leads to a tilt profile $\theta^{a,p}(x, y)$ which can be compared with $\theta_1(y)$ by considering the L_2 -norm

$$\begin{aligned} ||\theta^{a,p}(x, y) - \theta_1(y)|| = \\ \left(\int_0^1 \int_0^5 (\theta^{a,p}(x, y) - \theta_1(y))^2 dy dx \right)^{1/2}. \end{aligned} \quad (34)$$

From Figure 5, it is clear that it is far more efficient to make use of a coarse mesh and second- or third- order elements than it is to refine the mesh. For a given number of degrees of freedom, the higher p solution is typically an order of magnitude more accurate, and although not shown, the same is true with regard to the CPU time taken and the memory used.

3.2 Uniform h - and p -refinement close to a defect

At the other extreme of length scale, a test problem containing a defect was defined on a 50 nm square. Weak homeotropic anchoring with $W = 10^{-3} \text{Jm}^{-2}$ was imposed on the left-hand and lower walls, and natural boundary conditions on the remaining two. There is no analytic solution to this problem, so for the purpose of comparison we generated a reference solution, using $p = 2$ elements on a mesh of right angled triangles with two sides of length $h = 1$ nm. This reference solution, shown in Figure 6, is uniaxial nearly everywhere, except for a biaxial region in the few nanometres surrounding the bottom-left corner, where the defect is formed. Here we follow Barberi et al [9], in quantifying biaxiality by a measure b given by

$$b^2 = 1 - \frac{6\text{tr}(Q^2)^3}{\text{tr}(Q^3)^2} \quad (35)$$

which varies from $b = 0$ in uniaxial regions, to $b \leq 1$ in biaxial regions.

Making use of a mesh of $p = 2$, $h = 25$ nm triangles there is a marked departure from the reference solution. Although the director profile, plotted in Figure 7, away from the defect is much the same, the biaxial region is much larger, extending over a length scale of $h = 25$ nm - the size of one element. We can reduce the size of the biaxial region toward the correct value by either reducing the mesh spacing, as in Figure 8, where the use of $h = 6.25$ nm triangles results in a defect length scale close to the correct size, or by increasing the order, as in Figure 9, where $p = 6$ elements have been used to similar effect.

As in the previous test, the key question is one of efficiency, but the answer is less clear cut. If we set the number of degrees of freedom, then, again, a coarser mesh of higher order elements delivers the better accuracy, as can be seen in Figure 10. However, if we measure efficiency with respect to computational effort, as in Figure 11 then, while $p = 6$ elements outperform $p = 2$ elements, $p = 9$ elements perform similarly. In other words, the cost per degree of freedom increases more quickly as p increases than the error is reduced, so that at some point, p -refinement becomes the more costly, and as a result, we ought to choose mesh spacings so that $p \gg 9$ elements are not often needed. The primary cause is the spiralling complexity of (19) discussed earlier, although the resulting linear systems takes longer to solve as well.

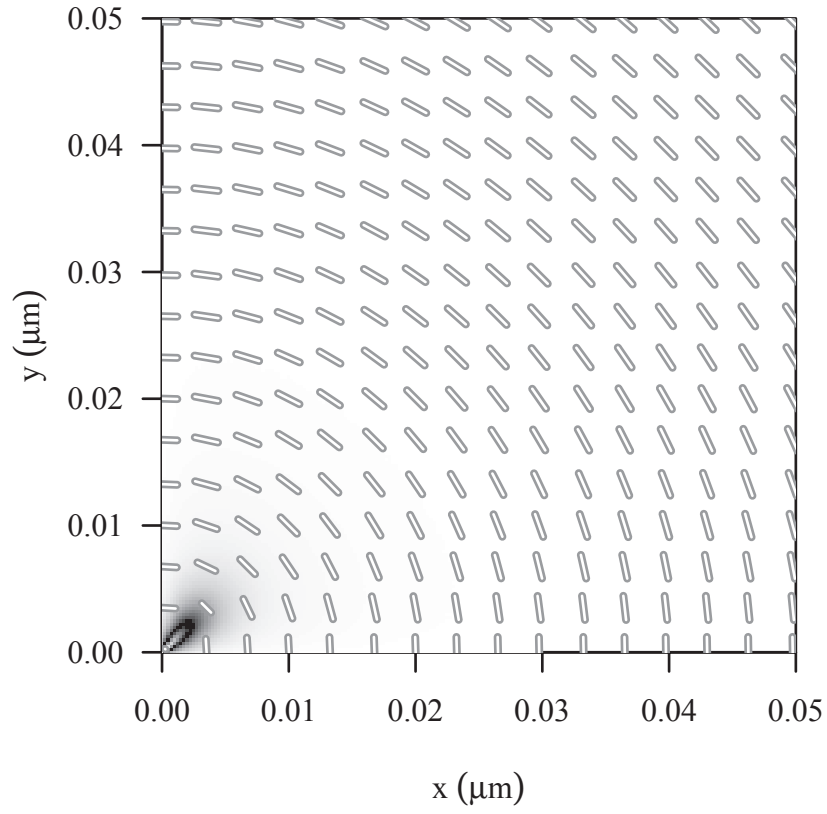


Figure 6: Reference solution for the corner problem. This solution is computed using second order elements on a fine, $h = 1$ nm mesh. The colour map depicts the biaxiality parameter, b , while the line segments are parallel to the director.

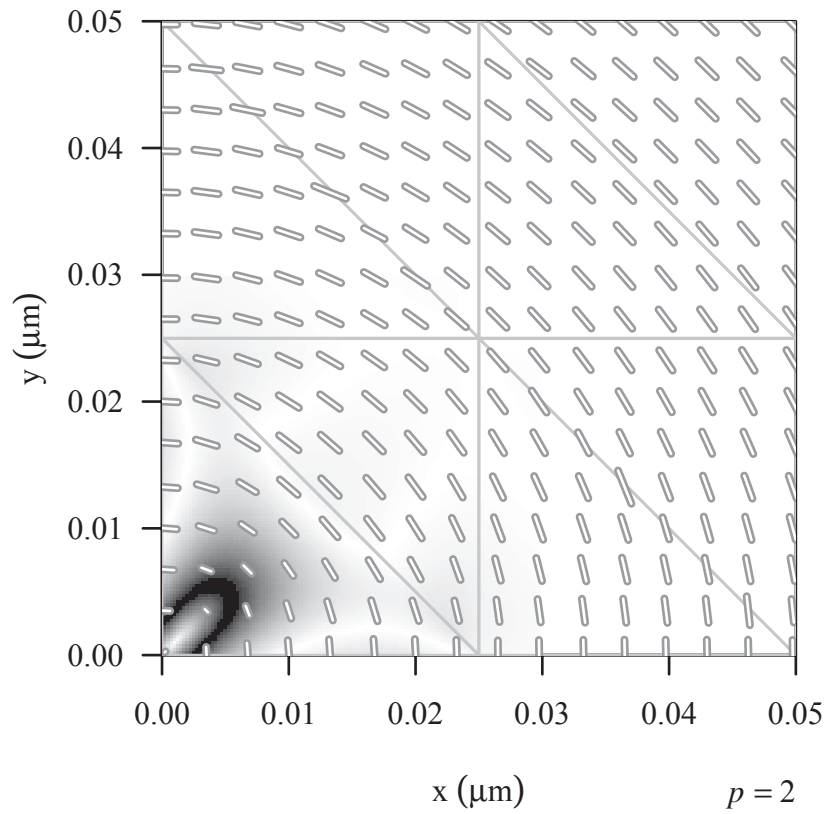


Figure 7: Solution for the corner problem using $h = 25$ nm, $p = 2$ triangular elements. In this solution, the biaxial region covers much of the bottom left element, and thus an exaggerated length scale of 25 nm.

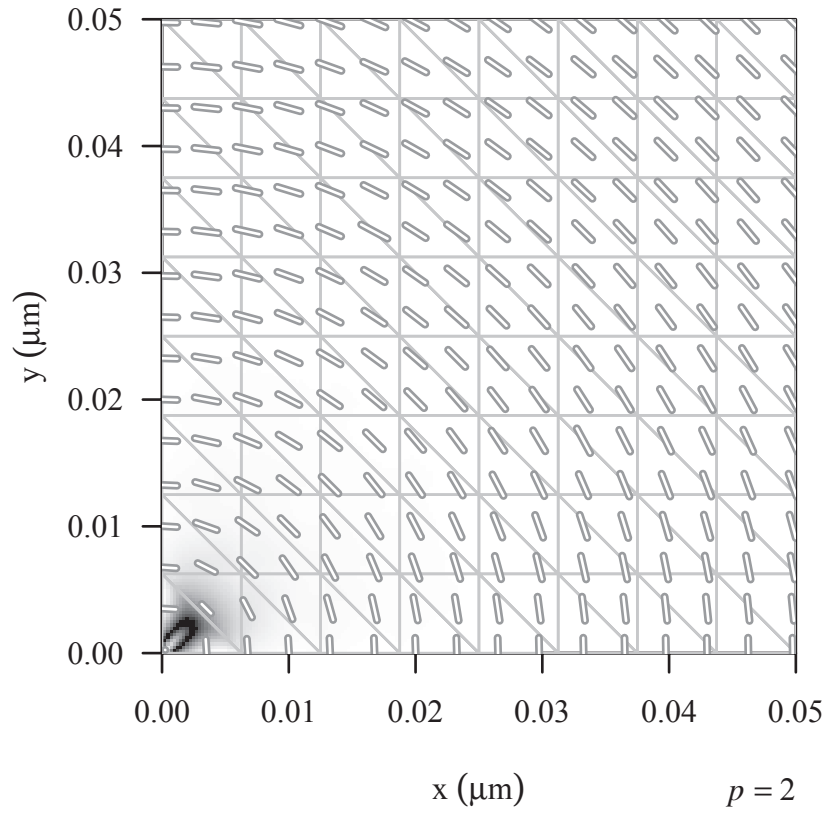


Figure 8: Solution for the corner problem using $h = 6.25$ nm, $p = 2$ elements. Here, the biaxial region fills all of the bottom left element and much of the adjacent element, but is now closer to the correct size because the elements have shrunk.

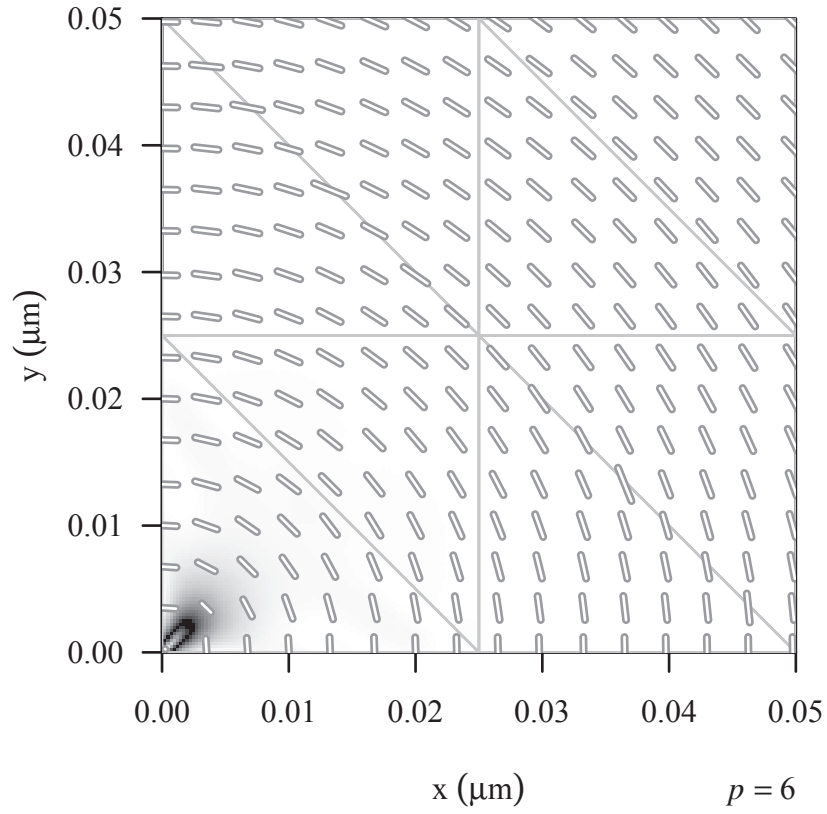


Figure 9: Solution for the corner problem using $h = 25$ nm, $p = 6$ elements. The biaxial region is close to the correct size, occupying only a fraction of the bottom left element.

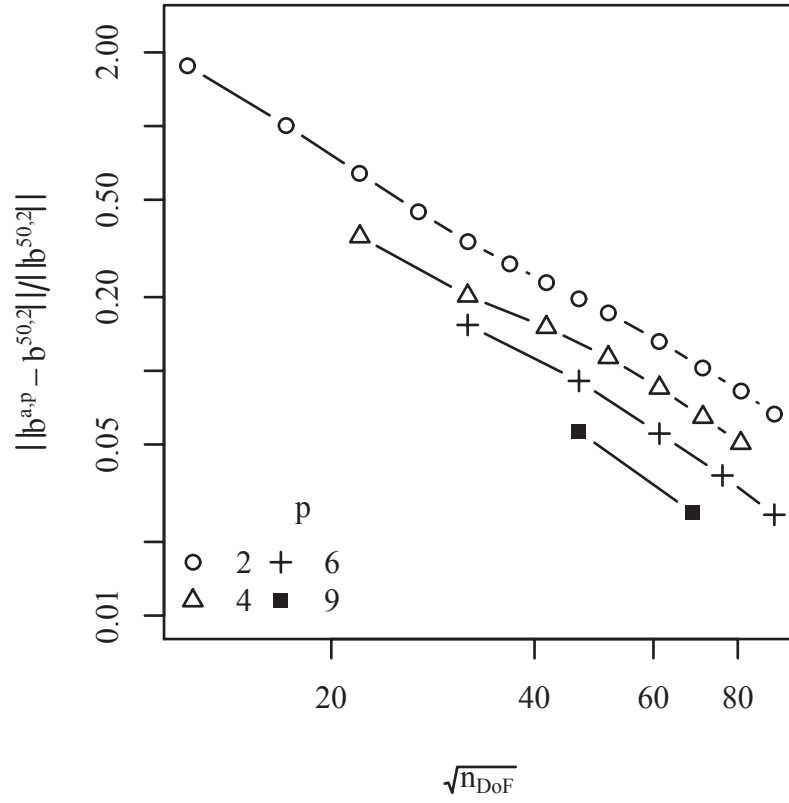


Figure 10: Error in the biaxiality plotted against number of degrees of freedom. Following any of the lines from left to right, the mesh spacing h is reduced while the polynomial order p is kept constant. For a given number of degrees of freedom (n_{DoF}), the higher p (and so lower h) solution is more accurate.

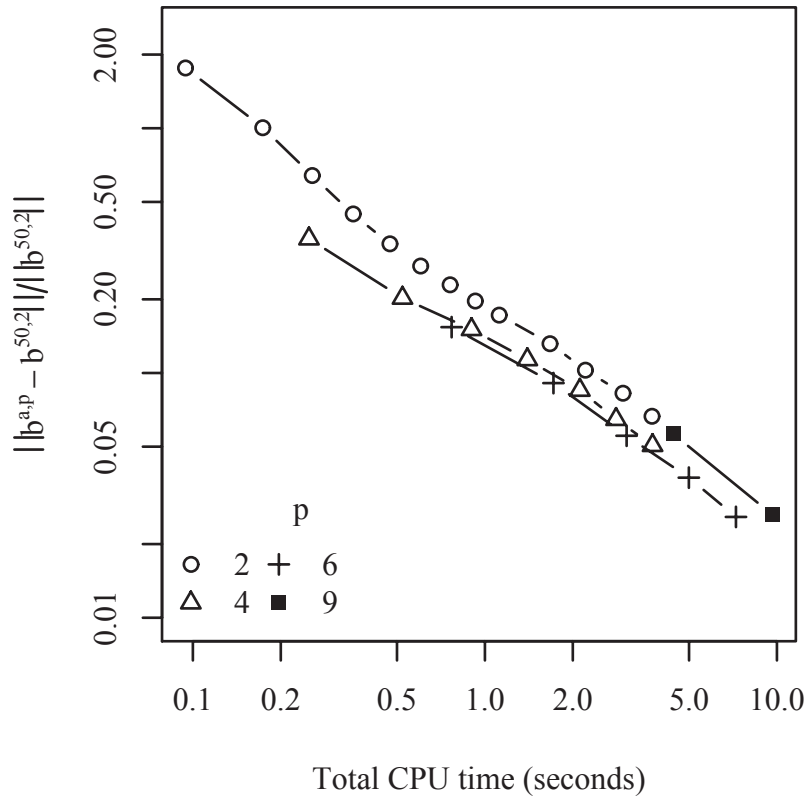


Figure 11: Error in the biaxiality plotted against number of degrees of freedom plotted against CPU time. When CPU time is considered rather than n_{DoF} , p -refinement loses some of its advantage. Nonetheless, for a given accuracy $p = 4$ or $p = 6$ elements require less time than $p = 2$ elements, and $p = 9$ elements require no more.

	h_b (nm)	n_{DoF}	Time (s)	Memory (Gb)	B (nm ²)
adaptive	15	41,370 [†]	301	1.21	16.3
fixed	15	38,358	212	0.82	35.1
	10	58,302	387	1.13	20.5
	8	77,226	420	1.44	18.5
	5	140,766	810	2.49	17.0
	3	270,582	1,404	4.63	16.6
	2	482,550	2,874	8.18	16.5

Table 1: Computational requirements compared with biaxial volume, B , for both the adaptive and fixed mesh methods. Not only is the adaptive method quicker to converge on a value of B around 16 nm², it uses less memory as well, even when the cost of error estimation is taken into account. [†]For the adaptive method, the number of degrees of freedom for the basic problem is given: 86,214 are needed for error estimation.

3.3 Automatic p -adaption applied to a realistic problem

We now turn our attention to a realistic problem, one in which both length scales need to be considered together. Figure 12 shows the geometry of the problem: a square box, 1.5 microns on a side, with rounded corners, Ω . The liquid crystal in bulk is uniaxial and aligned parallel to the boundary Γ all along its length. The mesh also shown in Figure 12 is the coarsest we will use, with the mesh spacing varying from $h_c = 80$ nm in the centre to $h_b = 15$ nm at the edges. This mesh, which we will refer to as R_{15} , is to be used with the automatic p -adaption techniques outlined earlier. For comparison, we also generated a sequence of meshes $R_{10}, \dots, R_2$ with exactly the same boundary as R_{15} , but with h_b decreasing down to 2 nm.

Although this problem does not have an analytic solution, its character is known, having been studied experimentally and theoretically [3]. All the steady state solutions contain two defects, with charge $c = +1/2$ - we will study the case where they sit close to the top left and bottom right corners. Figure 13 plots both the biaxiality b and the director over the whole of Ω , computed on mesh R_{15} with $p = 2$ elements. The director is illustrated well enough on this scale, and doesn't vary much between any of the computations, while the biaxiality varies only in a region close to the defects, barely visible on this scale.

Looking more closely at one of the defects in Figure 14, it is clear that some refinement is needed. Although the profile of b has the right character - uniaxial nearly everywhere, with a biaxial annulus centred on the defect - the biaxial region is somewhat triangular in shape, which we regard with suspicion as the finite elements are also triangular. Applying the automatic p -adaption technique, we arrive at the solution shown in Figure 15. A number of higher order elements have appeared, centred on the defect which sits inside a $p = 10$ element. The defect itself has taken on a circular shape, shrunk considerably, and sits inside a different element.

A similar improvement in the solution can be made by using a fixed mesh with finer elements close to the boundary, but the automatic p -adaption is far more efficient. In Figure 16, we plot, for both approaches, the computational time against the biaxial volume B , of the top-left defect,

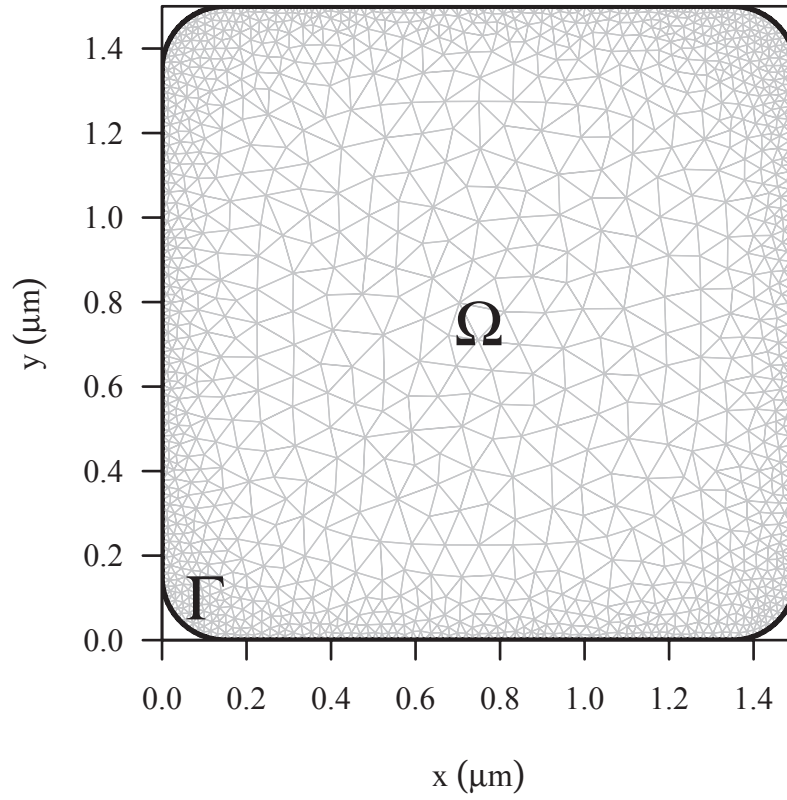


Figure 12: Mesh R_{15} for the rounded corner problem. The domain Ω is bounded by Γ , a square with rounded corners approximated by straight segments 15 nm long. Planar homogeneous anchoring, with the director in the plane of the figure, is imposed on the whole of Γ . In the centre of Ω , elements are $h_c \approx 80$ nm on a side, while at the boundary they are $h_b \approx 15$ nm on a side. The finer meshes R_2 - R_{10} have exactly the same boundary as R_{15} , but a lower value of h_b .

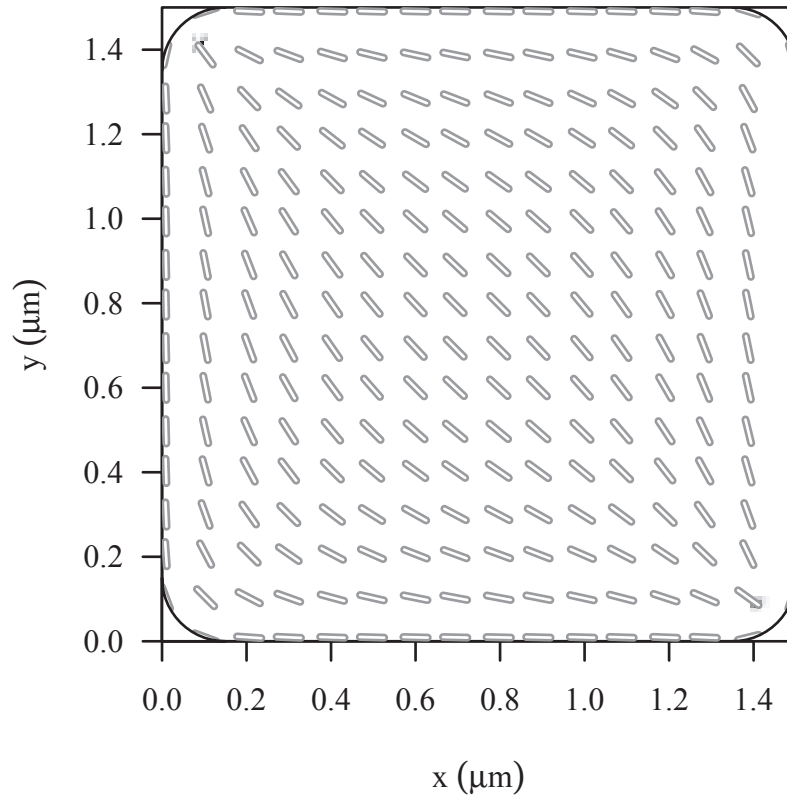


Figure 13: Biaxiality and director plot for the rounded corner problem. $c = +1/2$ defects can be seen close to top left and bottom right corners of the figure. Variation in the biaxiality is confined to a barely visible region.

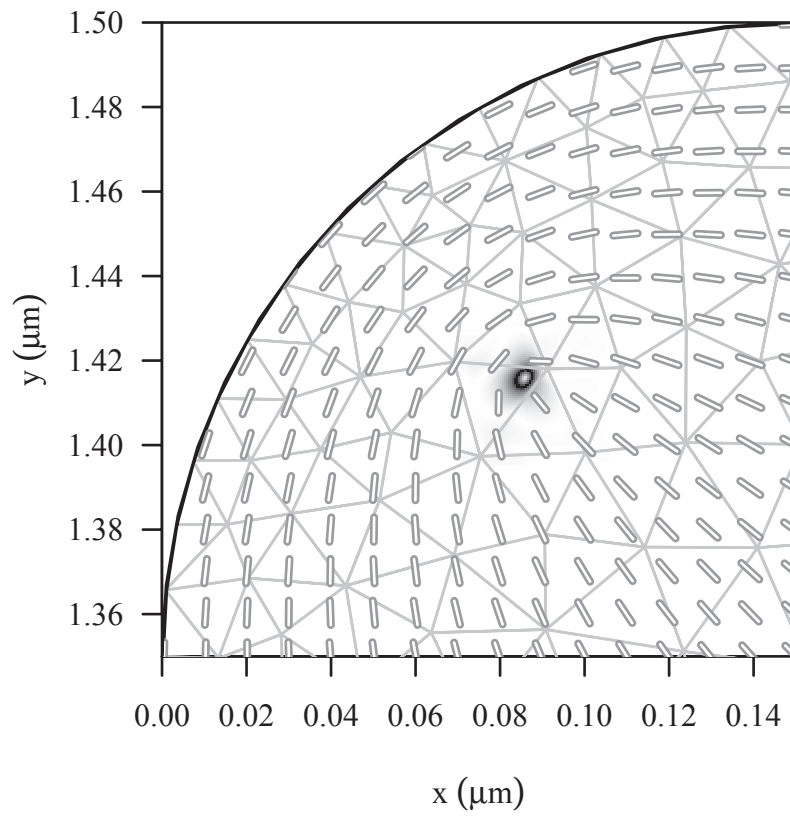


Figure 14: The top left corner of mesh R_{15} before any refinement. At this resolution the defect is more clearly visible, and has a triangular shape.

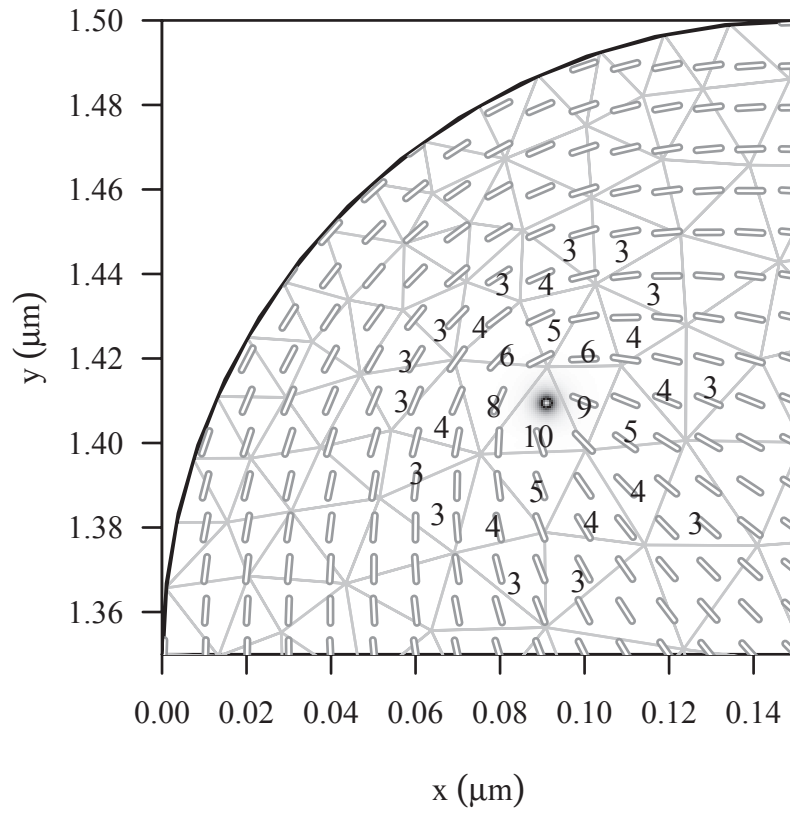


Figure 15: The top left corner of mesh R_{15} after the final refinement. There are now a number of (labelled) $p > 2$ order elements surrounding the defect, which has shrunk, become circular in shape, and moved position into an adjacent element.

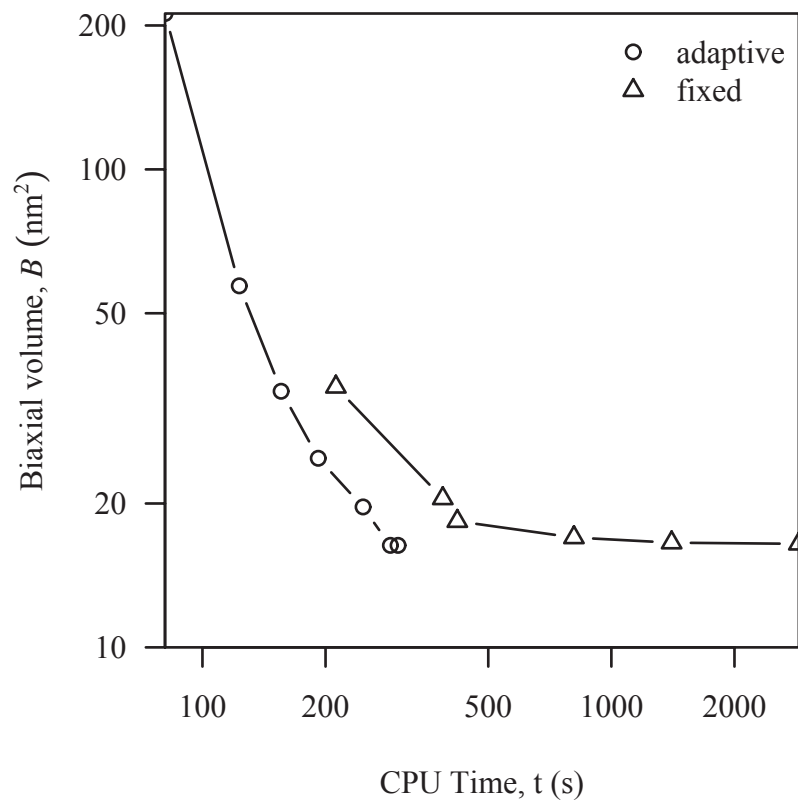


Figure 16: Biaxial volume plotted against CPU time for both adaptive and fixed mesh solutions. For a given volume, the adaptive method is faster.

found by integrating the biaxiality b over the portion of Ω shown in Figure 14 and Figure 15. For the automatic p -adaption, a value of $B = 16.3 \text{ nm}^2$ is reached after around 300 seconds, while it is necessary to use $h_b \leq 3 \text{ nm}$ and more than four times the computational effort to reach a similar result using a fixed mesh. These results are reiterated in Table 1, together with data showing that the p -refinement technique also uses far less memory, despite the additional storage needed to calculate the error estimates.

4 Conclusions

The work that we have presented here has been focussed on the modelling of realistically sized liquid crystal devices where, particularly when defects are involved, the different length scales in the problem make modelling very difficult.

As part of this investigation we have shown that in defect free regions using second, or third, order elements on a coarse mesh is far more efficient than using piecewise elements on a fine mesh.

The main focus of this work has been modelling regions containing defects and in this case we have shown that hierarchical finite elements can be used to account for the disparate length scales in the problem and that close to defects using elements of up to order nine is as efficient, or more efficient than, using lower order elements on a finer mesh.

Using these hierarchical elements we have implemented an adaptive technique which automatically discovers the location of defects and increases the order of the elements around them. The resulting method is simple to implement, and promises to be well suited to time-dependent problems, where the use of hierarchical function spaces will make interpolation of solutions from one mesh to the next straightforward.

In the future, we will report on the application of these methods to time-dependent and three-dimensional systems, and on progress in reducing the storage cost of the error estimation procedure.

Acknowledgements

This work was conducted at HP Laboratories, Bristol, UK as part of a joint EPSRC and HP Laboratories funded project. We are grateful to Professor Mark Ainsworth, Dr Alison Ramage, Dr Xinhui Ma and Professor Nigel Mottram, all at the University of Strathclyde, for their advice on both finite element methods and $\mathbf{Q}$ -tensor theory. GNU R [19] was used for data analysis.

References

- [1] Bryan-Brown G P, Brown C V, Jones, J C, Wood, E L, Sage I C, Brett P and Rudin J 1997 *SID Digest* **28** 37
- [2] Kitson S and Geisow A 2002 *Appl. Phys. Lett.* **80** 3635

- [3] Tsakonas C, Davidson A J, Brown C V, Mottram N J 2007 *Appl. Phys. Lett.* **90** 111913
- [4] Gartland Jr. E C, Palffy-Muhoray P and Varga R S 1991 *Mol. Cryst. Liq. Cryst.* **199** 429
- [5] Sonnet A, Kilian A and Hess S 1995 *Phys. Rev. E.* **52** 718
- [6] Ainsworth, M and Oden, J T *A posteriori error estimation in finite element analysis* Wiley-Interscience (2000)
- [7] Davidson, A J and Mottram, N J 2002 *Phys. Rev. E.* **65** 051710
- [8] Parry-Jones L A and Elston S J 2005 *J. App. Phys.* **97** 093515
- [9] Barberi R, Cuichi F, Durand G E, Iovane M, Sikharulidze D, Sonnet A M and Virga E G 2004 *Eur. Phys. J. E.* **13** 61
- [10] Willman E, Fernández F A, James R and Day S E 2008 *Journal of Display Technology* **4** 276
- [11] Grundmann, A and Möller, H M 1978 *SIAM. J. Numer. Anal.* **15** 282
- [12] Zienkiewicz O C and Zhu J Z 1991 *Int. J. Numer. Meth. Eng.* **32** 783
- [13] Coles, H 1978 *Mol. Cryst. Liq. Cryst.* **49** 67
- [14] Ainsworth M and Coyle J 2001 *Comput. Meth. Appl. Eng.* **190** 6709
- [15] Ainsworth M and Coyle J 2003 *Int. J. Numer. Meth. Eng.* **58** 2103
- [16] Zienkiewicz, O C and Taylor, R L *The finite element method for solid and structural mechanics* Butterworth-Hienemann (2000)
- [17] Balay S, Buschelman K, Eijkhout V, Gropp W D, Kaushik D, Knepley M G and McInnes L C 2008 *PETSc Users Manual* ANL-95/11 - Revision 3.0.0 (Argonne National Laboratory)
- [18] Amestoy P R, Duff I S, Koster J, L'Excellent J Y 2001 *SIAM J. Matrix. Anal. App.* **23** 15
- [19] R Development Core Team 2008 *R: A Language and Environment for Statistical Computing* (Vienna: R Foundation for Statistical Computing)