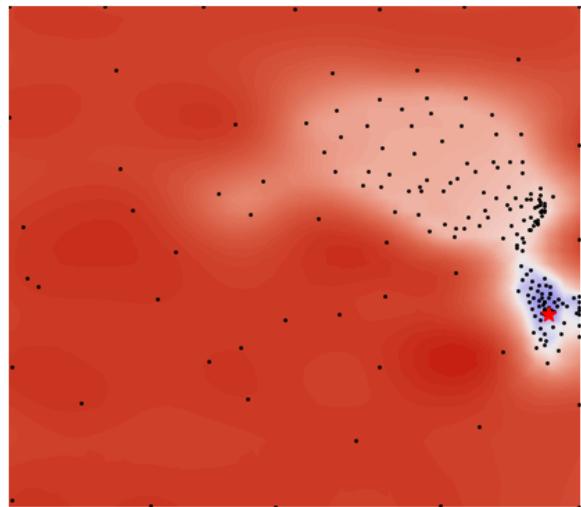
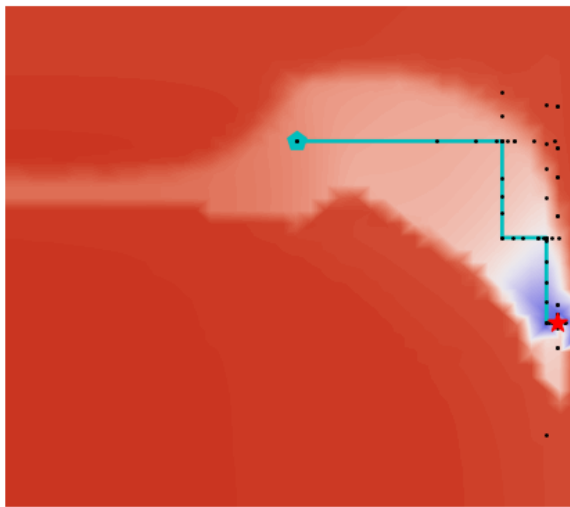




CHALMERS
UNIVERSITY OF TECHNOLOGY



Optimization of tokamak disruption scenarios

Avoidance of runaway electrons and excessive wall loads

Thesis for the degree of Master of Science in Physics

HANNES BERGSTRÖM
PETER HALLDESTAM

DEPARTMENT OF PHYSICS

CHALMERS UNIVERSITY OF TECHNOLOGY
Gothenburg, Sweden 2022

MASTER'S THESIS 2022

Optimization of tokamak disruption scenarios

Avoidance of runaway electrons and excessive wall loads

HANNES BERGSTRÖM
PETER HALLDESTAM



Department of Physics
Division of Subatomic, High Energy and Plasma Physics
CHALMERS UNIVERSITY OF TECHNOLOGY
Gothenburg, Sweden 2022

Optimization of tokamak disruption scenarios
Avoidance of runaway electrons and excessive wall loads
Hannes Bergström
Peter Halldestam

© HANNES BERGSTRÖM, PETER HALLDESTAM 2022.

Supervisors: István Pusztai, Department of Physics
Oskar Vallhagen, Department of Physics
Examiner: Tünde Fülöp, Department of Physics

Master's Thesis 2022
Department of Physics
Division of Subatomic, High Energy and Plasma Physics
Chalmers University of Technology
SE-412 96 Gothenburg
Telephone +46 31 772 1000

Cover: Two optimization methods, Powell's (left) compared to the Bayesian method (right). The left subfigure further illustrates a scan of 1600 fluid disruption simulations, from which a value of the objective function is evaluated, describing the performance of the disruption scenario. The trajectory taken by Powell's method is superimposed on top of the scan. Shown in the right subfigure, a point estimate of the objective function that is rooted in Bayesian statistics. For the complete figure, see Fig. 5.2.

Typeset in L^AT_EX
Printed by Chalmers Reproservice
Gothenburg, Sweden 2022

Hannes Bergström
Peter Haldestam
Department of Physics
Chalmers University of Technology

Abstract

Research in the field of fusion science has been propelled by its potential to alleviate humanity's reliance on fossil fuels. One of today's most promising approaches to generating thermonuclear fusion energy uses magnetic confinement of hydrogen fuel in the plasma state. The tokamak concept, which has achieved the best fusion performance so far, is used in the three reactor-scale devices (ITER, SPARC and STEP) currently being constructed — they aim to achieve a positive energy balance, thereby demonstrating the scientific feasibility of magnetic confinement fusion energy.

A major open issue threatening the success of these tokamaks is plasma disruption. In these off-normal events the plasma loses most of its thermal energy on a millisecond timescale, exposing the device to excessive mechanical stress and heat loads. In addition, in the high-current devices currently under construction, one of the most important related problems is posed by currents carried by electrons accelerated to relativistic energies, called runaway electrons. If these were to strike the inner wall unmitigated, it may cause potentially irreversible damage to the device. The methods proposed to mitigate these dangerous effects of disruptions, such as massive material injection, are characterized by a large number of parameters, such as when to inject material, in which form and composition. This poses an optimization problem that involves a potentially high-dimensional parameter space and a large number of disruption simulations.

In this work, we have developed an optimization framework that we apply to numerical disruption simulations of plasmas representative of ITER, aiming to find initial conditions for which large runaway beams and excessive wall loads can be avoided. We assess the performance of mitigation when inducing the disruption by massive material injection of neon and deuterium gas. The optimization metric takes into account the maximum runaway current, the transported fraction of the heat loss — affecting heat loads — and the temporal evolution of the ohmic plasma current — determining the forces acting on the device.

Keywords: fusion plasma, disruption mitigation, runaway electrons, Powell's method, Bayesian optimization.

Acknowledgements

We want to express our sincerest appreciation to everybody at the Plasma Theory group at Chalmers. First and foremost, we would like to thank our supervisors, István Pusztai and Oskar Vallhagen, for their tremendous help during our time with the group. They have always been available to answer any of our physics related questions and proofread our thesis. We are also very grateful for having Tünde Fülöp as our examiner and mentor. She has given us the amazing opportunity to work with this project, and her generous support and enthusiastic guidance has truly given us the ideal start of our scientific careers. We also want to thank Mathias Hoppe, the main developer of the DREAM-code. Finally, a special thanks to the computer scientists in the Energy Area of Advance project *Optimizing fusion with generative programming*, who have provided fruitful insights within the field of black-box optimization.

Hannes Bergström & Peter Haldestam, Gothenburg, June 2022

Contents

1	Introduction	1
1.1	Magnetic confinement fusion	2
1.2	Tokamak disruptions	2
1.3	Thesis outline	4
2	Review of plasma physics	5
2.1	Basic properties	6
2.2	Electromagnetic fields	8
2.3	Theoretical models	10
2.3.1	Kinetic theory	10
2.3.2	Fluid theory	11
2.4	Collisions	13
2.4.1	The Fokker-Planck collision operator	13
2.4.2	The runaway phenomenon	15
2.5	The tokamak	18
2.5.1	Fields and particle motion	19
2.5.2	Disruptions	20
2.5.2.1	Mitigation techniques	23
3	Disruption model	25
3.1	Current density evolution	26
3.2	Background plasma dynamics	26
3.3	Runaway electron generation	30
3.4	Numerical implementation	32
4	Black-box optimization	35
4.1	Powell’s conjugate direction method	36
4.1.1	The problem of linear dependence	38
4.1.2	Golden section search and Brent’s method	40
4.1.2.1	Bracketing a minimum	42
4.2	Bayesian optimization	43
4.2.1	Gaussian process regression	45
4.2.2	Expected improvement acquisition function	49
5	Optimization of disruption mitigation scenarios	51
5.1	Initially prescribed temperature evolution	52
5.2	Prescribed heat transport	56

5.3	Profile optimization	61
5.4	Sensitivity analysis	67
6	Conclusion and outlook	71
6.1	Conclusion	71
6.2	Outlook	73
	Bibliography	75
A	Analytic magnetic geometry	I
B	Gaussian process predictions	V
B.1	Initially prescribed temperature evolution	V
B.2	Pessimizations of found minima	VI

1

Introduction

Humanity's heavy reliance on burning fossil fuels and excessive emissions of greenhouse gases have caused global mean temperatures to increase more than 1°C in the last hundred years [1, 2]. While it is clear that the use of fossil fuels must be reduced as soon as possible, it is not as obvious what alternative sources of energy are to be their replacement. Renewable energy, such as wind and solar power, are able to produce electricity without emitting much greenhouse gases, but they are heavily dependent on local weather conditions, thus providing an unreliable supply of electricity to the grid [3].

Energy obtained from *nuclear fission*, used in modern power plants, provides the most power per unit area out of all energy sources currently available. Such nuclear reactions include neutron-induced fission of the uranium-235 isotope into various smaller nuclei. The difference in mass prior to and after the reaction is turned into kinetic energy for the reaction products. This energy can subsequently be used to power steam turbines that generate and provide electricity to the grid, making nuclear fission a much more predictable source of energy compared to the currently available renewable alternatives. However, concerns about the accumulation of long-lived radioactive waste and the threat of the extraordinary event of catastrophic nuclear meltdowns, such as those in Chernobyl (1986) and Fukushima (2011), all act to counterbalance the public opinion on whether or not to build more fission reactors. The world's energy dependence on fossil fuels is, along with the destructive potential of nuclear weapons, among the defining aspects of modern time.

Another category of nuclear reactions is *thermonuclear fusion*, in which nuclear energy is instead harnessed by fusing together smaller particles into more stable and heavier ones, rather than splitting them apart as in fission reactions. The extreme temperatures required for the kinetic energy of the nuclei to overcome the Coulomb barrier and fuse together makes it very difficult to achieve sustained nuclear fusion here on Earth. To obtain controlled fusion energy, one needs to introduce a sufficiently strong force field to confine the hot ionized gas constituting the fusion fuel. In stars, such as the Sun, a strong gravitational field puts the constituent *plasma* in force balance with the immense radiation pressures arising from the fusion reactions of hydrogen into helium in the core region.

The gargantuan mass needed, makes *gravitational confinement fusion* essentially only possible in stars. On Earth, there are however a few options to consider. Instead of the gravitational force, *inertial confinement fusion*, for instance, uses the inertia of a plasma to achieve thermonuclear fusion. This can be done by rapidly compressing a gas such that the fuel becomes sufficiently dense and hot for fusion reactions to occur, though only during a limited amount of time before the plasma

expands. The approach considered in this thesis confines the charged particles of the plasma with the use of magnetic fields.

1.1 Magnetic confinement fusion

Due to the extreme temperatures required for the hydrogen isotopes to fuse, fusion fuel used in *magnetic confinement fusion* (MCF) is inherently in the plasma state. We know from electrodynamics that charged particles, making up the ionized gas, move in gyro-orbits around magnetic field lines, as further detailed in Sec. 2.2. MCF uses the conductive property of plasmas to confine the fuel with a strong magnetic field. A naive approach is to simply shape the magnetic field into a torus, in an effort of keeping the charged particles in closed orbits. However, this would result in the fusion fuel drifting radially outwards, as further discussed in Sec. 2.1, and consequently, a loss of its confinement [4].

First conceptualized by Soviet physicists in the early 1950s, the tokamak¹ resolves this issue by introducing a poloidal component of the magnetic field by inducing a plasma current of several megaamperes to counteract such drifts. A schematic of this design is illustrated in Fig. 1.1. This plasma current allows the magnetic geometry to be azimuthally symmetric, yielding the exceptional energy confinement properties of the tokamak, while it is also a potential source of instability [6]. Major instabilities can lead to a rapid uncontrolled loss of plasma confinement, which is a key process in this thesis.

1.2 Tokamak disruptions

Arguably one of the largest threats to reliable tokamak operation are off-normal events known as *disruptions*, which are induced by a sudden loss of plasma confinement [7]. When this occurs, there is an ensuing heat and particle transport that results in a rapid temperature drop, accompanied by a decrease in the electrical conductivity of the plasma. The latter might strike as counter-intuitive since at lower temperatures, such as here on Earth, the resistivity in e.g. metals usually increases with temperature. In addition, due to the inductance of the system, the change in plasma current — that a reduction of conductivity would entail — is counteracted by the subsequent induction of an electric field. These effects have severe implications that can lead to intolerable damages to the tokamak device.

Firstly, the rapid release of thermal energy at the onset of the disruption can give rise to localized heat loads which can melt the plasma facing wall components in the reactor. Secondly, the reduced conductivity implies that the plasma current will eventually decrease as well. If this happens too quickly, it can induce substantial *eddy currents* in the surrounding structures. On the other hand, if the plasma current decays slowly, the loss of confinement can cause the plasma to drift towards the wall, and the current initially conducted in the plasma to start flowing

¹Transliteration of the Russian word Токамак, an acronym for тороидальная камера с аксиальным магнитным полем, which roughly translates to “toroidal chamber with an axial magnetic field” [5].

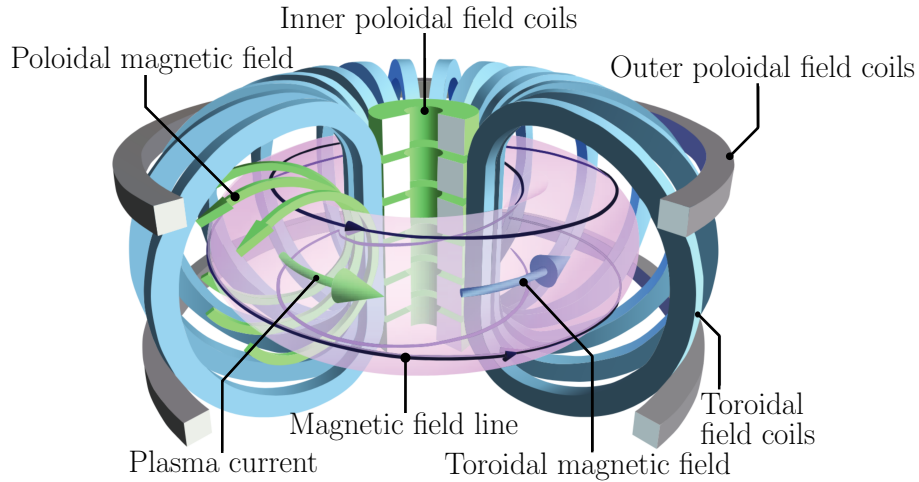


Figure 1.1: Schematic of the tokamak chamber and magnetic profiles. The poloidal and toroidal field components are highlighted, as well as the corresponding field coils. This setup yields magnetic field lines that wind around the torus in a helix. Source: [8].

in surrounding structures, a phenomenon referred to as a *halo current*. In both scenarios, the strong magnetic field permeating the device will cause forces to be exerted on the tokamak structure, potentially damaging the device. Finally, the induced electric field enables the existence of so called *runaway electrons* (REs), traveling at relativistic velocities. These highly energetic particles, if accumulated in sizable quantities, could potentially strike the wall with great intensity, leading to sub-surface melting of the components. The runaway phenomenon is described in more detail in Sec. 2.4.2.

To achieve dependable fusion energy using tokamak reactors, the mentioned effects need to be accounted for, either in the initial design of the device, or through the application of certain disruption mitigation techniques, introduced in Sec. 2.5.2.1. This problem is however quite complex, owing to the large number of different properties that factor in during disruption events, such as the geometry of the magnetic field, conducting properties of the wall, and initial temperature and current density profiles. Alternatively, the design of a viable disruption mitigation scheme can be viewed as an optimization problem, involving a large number of degrees of freedom. This is the topic of our thesis. The goal of the project is to construct a general optimization framework that can be applied to simulations of tokamak disruption scenarios, aiming to find a parameter regime in which the outcomes of the disruptions are tolerable. Since the simulations are based on solving highly non-linear systems of equations, implying that their behavior can be unpredictable, a black-box approach is adopted for the optimization. Two different optimization techniques, Powell's method and Bayesian optimization, are used, and their performance is evaluated.

1.3 Thesis outline

The chapters in this thesis are presented as follows; Ch. 2 gives a brief review of relevant aspects in the physics of plasmas for understanding tokamak disruptions. This is followed by a description, in Ch. 3, of the theoretical model used in this work for simulating disruptions in plasmas representative of the up-coming reactor-scale ITER² experiment. Next, Ch. 4 describes the black-box optimization algorithms, namely Powell’s method and Bayesian optimization, which are applied to various disruption simulations, as presented in Ch. 5. There, objective functions are constructed to account for the maximum runaway current, the transported fraction of the heat loss, as well as the current quench timescale. A range of multi-dimensional optimizations of disruption scenarios are carried out, particularly with regards to the injection of deuterium and neon gas. The conclusions are summarized in Ch. 6, including a discussion on further possible improvements regarding the optimization routines, and additional problems worth addressing with the optimization framework.

²ITER is a reactor size fusion experiment with the aim to exceed breakeven, currently built in Cadarache, France, in an international collaboration. Formerly the International Thermonuclear Experimental Reactor (ITER), it is now officially named for its Latin meaning: “the way”.

2

Review of plasma physics

When heating up a gas to sufficiently high temperatures, the individual atoms in it begin to lose their bound electrons, leaving a sea of both charged and neutral particles. The gas transitions into an ionized gas, or *plasma*. It is commonly referred to as the fourth fundamental state of matter, along with the solid, liquid and gaseous states. It has similarities to a gas but its ionization degree is sufficiently high that its dynamics are dominated by long-range electromagnetic interactions. In fact, only a small degree of ionization is required for it to exhibit electromagnetic properties. For instance, the electrical conductivity of the plasma already reaches half its maximum at a degree of 1% ionization [9].

Plasmas are abundant in our universe, existing across several orders of magnitude in both temperature and pressure. The Sun and other stars are massive luminous spheres of hot plasma confined by gravity, with their stellar winds continuously feeding their interstellar environment by a plasma outflow. Indeed, a vast majority of ordinary matter in the observable universe is in the plasma state [10]. Earth and its lower atmosphere, however, is an exception due to its relatively low temperatures. Interestingly, plasmas play an important role in e.g. electrical discharges during thunderstorms [11] and auroras generated by disturbances in Earth's magnetosphere caused by the solar wind [12].

With the aim to better understand the underlying physics of tokamak disruptions, we will in this chapter review a few core concepts within the field of plasma physics. Sec. 2.1 gives a brief summary of its relevant characteristics, followed by a description in Sec. 2.2 of the dynamics of charged particles in electromagnetic fields, particularly in the context of magnetically confined plasmas. In Sec. 2.3, we introduce the statistical framework for modeling the physics of plasmas used in this work. Next, Sec. 2.4 describes the collisional interactions between charged particles in a plasma, from which the phenomenon of runaway electrons emerges. Finally, in Sec. 2.5, we present a concise description of one of the core topics in this thesis, namely the tokamak and its outstanding issue of disruptions.

2.1 Basic properties

Formally, there is no exact definition for when a set of particles are said to have transitioned into the plasma state [13]. While it is achieved through ionization, a plasma can be either partially- or fully ionized, the extent of which is quantified by the *ionization degree* α , defined as

$$\alpha := \frac{n_i}{n_n + n_i}. \quad (2.1)$$

Here n_i and n_n denote the number density of ions and neutrals respectively. Note that an ionization degree of $\alpha > 1\%$ is often sufficient for a gas to behave like a plasma. The ionization state during heating of a weakly ionized plasma can be estimated using the *Saha ionization equation*

$$\left(\frac{n_i}{n_n}\right)^2 = 2.4 \cdot 10^{15} \frac{T^{3/2}}{n_n} e^{-U_i/k_B T}, \quad (2.2)$$

which assumes that the gas is in local thermal equilibrium at temperature T [6]. The form given in Eq. (2.2) specifically considers a gas consisting of a single atomic species with ionization energy U_i . For simplicity, from now on we make the Boltzmann constant k_B implicit by denoting the temperature in units of energy such that $k_B T \rightarrow T$.

Despite the ambiguity regarding phase transitions, there are some characteristic plasma properties that can be used to make the concept more explicit. A plasma can be described as a quasi-neutral gas of charged and neutral particles, which exhibit a collective behavior [6]. The term quasi-neutral implies that, while individual free particles may carry charge, the integrated charge density is approximately zero at sufficiently large length scales. Furthermore, should any charge imbalance arise, a plasma is able to maintain neutrality through a subsequent flow of particles that screen the induced Coulomb potential. To quantify this effect, consider a scenario where a test charge, q , is inserted into a plasma with spatially uniform electron number density, n_0 , and temperature, T_e , giving rise to a Coulomb potential $\phi(r)$. In addition, assume that the ions remain stationary due to their comparably large mass, while the perturbed electron density follows a Boltzmann distribution $n_e(r) = n_0 \exp(e\phi/T_e)$, with e denoting the elementary charge. It can then be shown that the resulting potential is given by

$$\phi(r) = \frac{q}{4\pi\epsilon_0 r} e^{-r/\lambda_D}. \quad (2.3)$$

Here $\lambda_D = \sqrt{\epsilon_0 T_e / n_0 e^2}$ is known as the *Debye length* and it characterizes the length scale of the screening effect. The assumptions used to derive Eq. (2.3) require that a large number of electrons are confined within this distance of the test charge. In other words, it should hold that $4\pi n_0 \lambda_D^3 / 3 \gg 1$. Another property that distinguishes a plasma is the collisionality. Collisions that occur between the charged particles in a plasma differ fundamentally from those between particles in a neutral gas because of the long range of the Coulomb force. Neutral particles collide in a manner reminiscent of a game of billiards, where a particle (or ball) in motion will interact

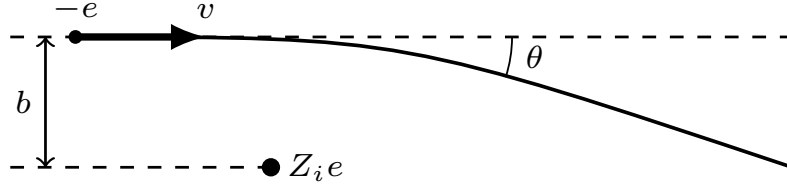


Figure 2.1: Illustration of the Coulomb collision between an electron with initial speed v and a stationary ion with charge $Z_i e$. Also shown is the impact parameter b and scattering angle θ .

with the system through close collisions, subsequently scattering with sharp deflection angles. Charged particles traveling through a plasma will on the other hand be repelled/attracted at larger distances, rather than colliding head-on, resulting in a series of *small-angle collisions*. Such an elementary interaction is illustrated in Fig. 2.1, where a negatively charged electron passes by a positively charged ion, and is deflected by an angle θ from the ensuing attraction. Also shown in the figure is the *impact parameter*, b , describing the perpendicular distance between the electron trajectory and the center of the Coulomb potential generated by the ion. Deriving the frequency of collisions involves integrating over this parameter, for which a lower- and an upper cut-off are needed to avoid divergences at both ends. As an upper cut-off, it is appropriate to use the Debye length λ_D , since Coulomb interactions beyond this distance are screened. The lower limit is set to the *distance of closest approach*, denoted $b_{\min}$, which describes the shortest possible distance between two colliding particles, and can be derived through kinematics. Note that both of these quantities are dependent on the energy of the particles. In computing the integral, one obtains a factor, $\ln \Lambda$, known as the *Coulomb logarithm*, where $\Lambda := \langle \lambda_D / b_{\min} \rangle_M$ is the ratio between the Debye length and distance of closest approach averaged over a Maxwellian velocity distribution. If $\ln \Lambda \gg 1$, then the distance of closest approach is significantly smaller than the screening limit, implying that for most interactions $b \gg b_{\min}$, which corresponds to a less direct collision. In this scenario, small-angle collisions will therefore dominate. The Coulomb logarithm accounts for the cumulative effect of small-angle collisions when deriving the collision frequency. If we assume that a single collision corresponds to a change in particle momentum equal to $\Delta p \sim mv$, it can be shown that the frequency of collisions for particle species a , with particles species b , is given by

$$\hat{\nu}_{ab} = \frac{n_b q_a^2 q_b^2 \ln \Lambda}{4\pi \epsilon_0^2 m_a^2 v_{\text{th}}^3}. \quad (2.4)$$

Here n_b and q_b denote the density and charge of particle species b , respectively, while q_a , m_a and $v_{\text{th}} = \sqrt{2T_a/m_a}$ are the respective charge, mass, and thermal velocity of particle species a [14]. An important result from this is that, since the conductivity, σ , in a plasma is inversely proportional to the electron-ion collision frequency $\hat{\nu}_{ei}$, it has a temperature dependence according to $\sigma \propto T_e^{3/2}$.

2.2 Electromagnetic fields

Since a plasma is composed of free charged particles, its dynamics is governed by electromagnetic fields. Particle motion is then dictated by the *Lorentz force*

$$\mathbf{F}_L := q[\mathbf{E} + \mathbf{v} \times \mathbf{B}], \quad (2.5)$$

where $\mathbf{E}$ and $\mathbf{B}$ denote the electric- and magnetic fields and $\mathbf{v}$ is the particle velocity. The electromagnetic fields are in turn described by *Maxwell's equations*

$$\nabla \cdot \mathbf{E} = \frac{\rho}{\varepsilon_0}, \quad (2.6a)$$

$$\nabla \cdot \mathbf{B} = 0, \quad (2.6b)$$

$$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t}, \quad (2.6c)$$

$$\nabla \times \mathbf{B} = \mu_0 \left(\mathbf{j} + \varepsilon_0 \frac{\partial \mathbf{E}}{\partial t} \right). \quad (2.6d)$$

The first term in Eq. (2.5) describes an acceleration along the electric field, while the second term is a bit more involved. Consider a scenario where there is a uniform magnetic field pointing in the z -direction, that is, $\mathbf{B} = (0, 0, B)$. Furthermore, assume that there is no electric field, i.e. $\mathbf{E} = 0$, and that a particle with charge q and mass m is moving in this region with velocity $\mathbf{v} = (v_x, v_y, v_z)$. The equations of motion for the particle is given by

$$\dot{\mathbf{x}} = \mathbf{v}, \quad \dot{\mathbf{v}} = \frac{q}{m} \mathbf{v} \times \mathbf{B}. \quad (2.7)$$

Expanding the different components in the acceleration $\dot{\mathbf{v}}$, one finds

$$\dot{v}_x = \frac{qB}{m} v_y, \quad \dot{v}_y = -\frac{qB}{m} v_x, \quad \dot{v}_z = 0. \quad (2.8)$$

Note that there is no acceleration in the direction parallel to the magnetic field, while the derivative is non-zero in the orthogonal xy -plane. Furthermore, taking the time derivative of Eq. (2.8) and inserting the expressions for $\dot{v}_x$ and $\dot{v}_y$, yields

$$\ddot{v}_x = -\left(\frac{qB}{m}\right)^2 v_x, \quad \ddot{v}_y = -\left(\frac{qB}{m}\right)^2 v_y. \quad (2.9)$$

Letting $\mathbf{v}_\perp := -\mathbf{B} \times (\mathbf{B} \times \mathbf{v})/B^2$ denote the velocity component perpendicular to the magnetic field lines, we may write $\ddot{\mathbf{v}}_\perp = -\omega_c^2 \mathbf{v}_\perp$.

This is the equation for a harmonic oscillator with angular frequency $\omega_c := |q|B/m$, implying that the electron gyrates around the magnetic field lines in a helical (circular + translational) orbit, as illustrated in Fig. 2.2. The radius of this orbit is given by $r_L = v_\perp/\omega_c$. The quantities ω_c and r_L are known as the *cyclotron frequency* and *Larmor radius* respectively, and they are fundamental properties for particle motion under the influence of magnetic fields. Note that the sign of q in Eq. (2.8) will determine the direction in which the particle is accelerated. As a result, positively charged particles rotate clockwise with respect to the magnetic field lines, whereas negatively charged particles rotate counterclockwise. Finally, the closed

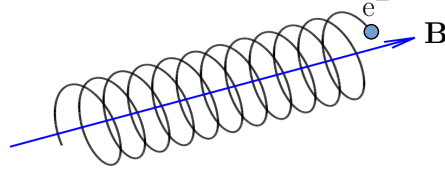


Figure 2.2: Electron in a gyro-orbit around a magnetic field line.

current loop resulting from the gyro-motion generates a magnetic moment

$$\mu := \frac{mv_{\perp}^2}{2B}. \quad (2.10)$$

An important property of μ is that it can be shown to be an *adiabatic invariant*, implying that it is approximately constant when the magnetic field has a slow spatio-temporal variation (compared to ω_c and r_L scales) [9]. The circular orbit entails greater complexity for particle trajectories. When studying more large scale phenomena it is therefore common to use approximations in the form of *drift-kinetic theory*. The model assumes that $r_L/L_B \ll 1$ and $\omega_B/\omega_c \ll 1$, where L_B and ω_B^{-1} are the respective characteristic length- and time scales for the magnetic field. Using perturbation theory, the perpendicular velocity of a particle can then be approximated as

$$\mathbf{v}_{\perp}(\mathbf{x}, t) \approx \mathbf{v}_D(\mathbf{x}) + \mathbf{v}_L(\mathbf{x}, t). \quad (2.11)$$

Here, $\mathbf{v}_L$ describes gyro-motion around the magnetic field lines, while $\mathbf{v}_D$ represents some drift velocity that varies on a larger time- and length scale. There are several mechanisms that can induce particle drifts, two of which have great significance in magnetic confinement fusion.

To begin with, if the magnetic field is spatially varying (as it must be to confine particles in a finite volume), the forces acting on a particle during a single gyration do not cancel out, resulting in a net velocity. Based on Eq. (2.11), it can be shown that the arising drift, $\mathbf{v}_{\nabla B}$, is given by

$$\mathbf{v}_{\nabla B} := \frac{\mu}{q} \frac{\mathbf{B} \times \nabla B}{B^2}, \quad (2.12)$$

where $B(\mathbf{x}) := |\mathbf{B}(\mathbf{x})|$ is the magnitude of the magnetic field. It should be pointed out that $\mathbf{v}_{\nabla B}$ will have opposite directions for positively- and negatively charged particles. Furthermore, an electric field would similarly accelerate and decelerate the particles along the gyration, yielding a drift velocity, $\mathbf{v}_{E \times B}$, described by

$$\mathbf{v}_{E \times B} := \frac{\mathbf{E} \times \mathbf{B}}{B^2}. \quad (2.13)$$

Unlike Eq. (2.12), this drift is independent of particle charge (this is the velocity of the unique frame in which the electric field Lorentz transforms to zero).

2.3 Theoretical models

Consider a plasma consisting of N particles. One possible approach to modeling this plasma would be to simply employ the equations of motion for each particle

$$\dot{\mathbf{x}}_i = \mathbf{v}_i, \quad \dot{\mathbf{v}}_i = \frac{q_i}{m_i}(\mathbf{E} + \mathbf{v}_i \times \mathbf{B}), \quad (2.14)$$

alongside Maxwell's equations for the electromagnetic field as given in Eq. (2.6). This however results in a set of $6N + 8$ coupled integro-differential equations, which quickly becomes numerically unfeasible to solve for systems of many particles. The complexity necessitates the introduction of reduced models. In the field of plasma physics these models are typically divided into two categories: kinetic theories and fluid theories. In this section we provide a brief overview of the fundamental equations and assumptions that are inherent to kinetic theory and fluid theory. For a complete derivation of the models the reader is referred to Ref. [9].

2.3.1 Kinetic theory

Let $\mathbf{z}_i = (\mathbf{x}_i, \mathbf{v}_i)$ denote the phase space coordinates of particle i . A system of N particles can then be fully defined by its N -particle distribution function

$$f_N(\mathbf{z}, \mathbf{z}_1, \dots, \mathbf{z}_N, t) := \prod_{i=1}^N \delta(\mathbf{z} - \mathbf{z}_i(t)). \quad (2.15)$$

Taking the time derivative of Eq. (2.15) one obtains the *Liouville equation*:

$$\frac{df_N}{dt} := \frac{\partial f_N}{\partial t} + \mathbf{v} \cdot \frac{\partial f_N}{\partial \mathbf{x}} + \mathbf{a} \cdot \frac{\partial f_N}{\partial \mathbf{v}} = 0, \quad (2.16)$$

where $\mathbf{a} = \partial \mathbf{v} / \partial t$. The equation implies that the distribution function remains constant along the phase space trajectories of the system.

While Eq. (2.16) may be a compact way of expressing the equations of motion, the degrees of freedom have not yet been reduced. To lower the dimensionality of the distribution function one can use the Bogoliubov–Born–Green–Kirkwood–Yvon (BBGKY) hierarchy which formulates a relation between the n -particle and $n + 1$ -particle distribution functions, together with cluster expansion [9]. There are two underlying assumptions for this step: first the BBGKY hierarchy requires a closure assumption regarding f_{n+1} , and secondly, it is assumed during the cluster expansion that the particles are fully uncorrelated. Defining a renormalised one-particle distribution function $f(\mathbf{x}, \mathbf{v}, t) := N f_1(\mathbf{x}, \mathbf{v}, t)$, one obtains a reduced kinetic equation from Eq. (2.16), known as the *Vlasov equation*. It assumes that the distribution function is constant along a given phase space trajectory, and it does not consider effects such as collisions between particles (two-particle correlations). This can be taken into account via the introduction of a collision operator, $C\{f\}$, yielding

$$\frac{\partial f}{\partial t} + \mathbf{v} \cdot \frac{\partial f}{\partial \mathbf{x}} + \mathbf{a} \cdot \frac{\partial f}{\partial \mathbf{v}} = C\{f\}, \quad (2.17)$$

which is called the *Boltzmann equation*. The exact form of the collision operator may

vary in several ways, for example if relativistic effects are included or only certain particle species are considered. It is however always constrained by fundamental conservation laws.

According to the H-theorem, the entropy of the system, $S := -k_B \int dv^3 f \ln(f)$, must follow

$$\frac{dS}{dt} \geq 0. \quad (2.18)$$

If one assumes an equilibrium state where $dS/dt = 0$, one can derive a solution to Eq. (2.17) called the *Maxwell-Boltzmann distribution*:

$$f_M(\mathbf{x}, \mathbf{v}, t) := n(\mathbf{x}, t) \left(\frac{m}{2\pi T(\mathbf{x}, t)} \right)^{3/2} \exp\left(-\frac{m(\mathbf{v} - \mathbf{u}(\mathbf{x}, t))^2}{2T(\mathbf{x}, t)} \right). \quad (2.19)$$

Here $n(\mathbf{x}, t) := \int dv^3 f$ is the particle density and $\mathbf{u}(\mathbf{x}, t) := \int dv^3 f \mathbf{v} / n$ is the mean velocity. These two properties are sometimes referred to as fluid quantities, a concept that allows us to reduce the complexity further, through what is known as fluid theory.

2.3.2 Fluid theory

In the previous section we were able to decrease the number of degrees of freedom considerably using kinetic theory. It is however possible to go even further and turn the 7-dimensional description into a 4-dimensional one (i.e. moving from $\{\mathbf{x}, \mathbf{v}, t\}$ to $\{\mathbf{x}, t\}$), by averaging over the velocity space. Such reduction of models is called fluid theory, and the associated quantities are derived by taking the *fluid moment* of the distribution function, defined as

$$\|Y\| := \int d^3v f(\mathbf{x}, \mathbf{v}, t) Y. \quad (2.20)$$

We have already seen examples of such moments in the form of particle density and mean velocity, corresponding to $Y = 1$ and $Y = \mathbf{v}/n$. A more comprehensive list of relevant fluid moments is given in Tab. 2.1. Note that the quantities depend on different powers of the velocity. They are therefore categorized into different orders, such that for example n is of the 0th order while $\mathbf{u}$ is of the 1st order.

In fluid theory, the governing equations can be derived by taking the moment of the kinetic equations presented in Sec. 2.3.1. It can for example be shown that taking the $Y = 1$ moment of the Boltzmann equation in Eq. (2.17) yields

$$\frac{\partial n}{\partial t} + \nabla \cdot (n\mathbf{u}) = 0. \quad (2.21)$$

This is known as the *continuity equation* and it describes conservation of particles. Furthermore, taking higher order moments one obtains the *momentum density equation* ($Y = m\mathbf{v}$) and the *energy density equation* ($Y = m\mathbf{v}^2/2$) given, respectively, by

$$\frac{\partial}{\partial t}(mn\mathbf{u}) + \nabla \cdot (mn\mathbf{u}\mathbf{u}) = -\nabla \cdot \mathbb{P} + nq(\mathbf{E} + \mathbf{u} \times \mathbf{B}) + \int d^3v C\{f\}\mathbf{v}, \quad (2.22)$$

$$\frac{\partial w}{\partial t} + \nabla \cdot \mathbf{Q} = qn\mathbf{u} \cdot \mathbf{E} + \int d^3v C\{f\} \frac{mv^2}{2}. \quad (2.23)$$

For later use it should be noted that the motion of charged particles can be expressed as a current density $\mathbf{j} = qn\mathbf{u}$, while the energy flux density can, under certain conditions, be written in an advection-diffusion form $\mathbf{Q} = -\mathbf{A}w + \mathbb{D}\nabla w$, where $\mathbf{A}$ is the advection vector and $\mathbb{D}$ is the diffusion tensor.

A crucial observation is the fact that each new equation contains an unknown fluid quantity of higher order. For example Eq. (2.21) depends on both the density n and mean velocity $\mathbf{u}$, while Eq. (2.22) also depends on the pressure tensor $\mathbb{P}$ and Eq. (2.23) in turn introduces the energy flux density $\mathbf{Q}$. This trend continues for higher order moments as well and it is therefore not possible to solve the system of equations based on the moments alone. An additional assumption is required in order to close the system. This is known as *fluid closure*.

There are many different closure strategies that range from simple constraints, such as the temperature being constant, to more complex models. In this work the closure is done by prescribing the particle and heat transport represented by the $\nabla \cdot (n\mathbf{u})$ term in Eq. (2.21) and the $\nabla \cdot \mathbf{Q}$ term in Eq. (2.23) respectively. A more exact description of how this is done is given in Sec. 3.2 and 3.3.

Table 2.1: List of relevant fluid moments, with the central velocity $\tilde{\mathbf{v}} := \mathbf{v} - \mathbf{u}$.

Fluid quantity	Definition
Particle density	$n := \ 1\ $
Mean velocity	$\mathbf{u} := \ \mathbf{v}\ /n$
Pressure tensor	$\mathbb{P} := m\ \tilde{\mathbf{v}}\tilde{\mathbf{v}}\ $
Energy density	$w := m\ v^2\ /2$
Energy flux density	$\mathbf{Q} := m\ v^2\mathbf{v}\ /2$

Furthermore, there is an additional reduction often used when studying evolution of conducting fluids and plasmas, called *magnetohydrodynamics* (MHD) [15]. It follows the fluid approach by combining the equations of motion for a given moment of each fluid species into a single one. In this description of a plasma, center of mass quantities are instead considered, as well as their coupling to the electromagnetic fields. The characteristic velocity appearing in dynamic solutions of the MHD system is the Alfvén speed. In order to meaningfully talk about magnetic confinement, it is clear that the plasma cannot evolve on such time scales, and as such it should be in a stable MHD equilibrium, given by the force balance

$$\nabla P = \mathbf{j} \times \mathbf{B}, \quad (2.24)$$

together with Eq. (2.6b) and Ampère’s law (2.6d) (neglecting the time derivative of the electric field). In Eq. (2.24), P is the isotropic pressure of the combined fluid.

2.4 Collisions

As previously mentioned in Sec. 2.1, in a plasma with a large Coulomb logarithm, the dominant type of collisions are small-angle collisions. In such a collision the velocity vector of an incoming particle is deflected only by a small amount, implying that the velocity changes gradually over time. Particles follow a more smooth trajectory compared to an ordinary gas, in which particles are colliding elastically with each other making them change direction very rapidly. Using the assumption that $\ln \Lambda \gg 1$ makes it possible to model collisions in a plasma with the so-called *Fokker-Planck collision operator*. In the following section we consider the non-relativistic Fokker-Planck collision operator in order to understand the underlying physics regarding the runaway phenomenon.

2.4.1 The Fokker-Planck collision operator

With the right-hand side set to zero, Eq. (2.17) would describe a flow of f in phase space such that the number of particles is conserved in any infinitesimal phase space volume advected with this flow. Collisions, described by the right-hand side of Eq. (2.17), can however scatter particles in and out of such volumes, while still conserving the number of particles globally.

Consider a particle in one dimension that is described by a distribution function $f(v, t)$. Note that collisions change the momentum of particles but not their location, which is why the explicit spatial dependence is omitted from f . Let $F(v, \Delta v)$ denote the probability that a particle will change its velocity v by Δv during a short time Δt due to collisions. The distribution function will evolve during this time as

$$f(v, t + \Delta t) = \int d\Delta v f(v - \Delta v, t) F(v - \Delta v, \Delta v). \quad (2.25)$$

In this expression the integrand can be interpreted as the probability density of particles with an initial velocity $v - \Delta v$ attaining a velocity v during this time. Integrating over all such changes in velocity will therefore yield the total density of particles with velocity v at the later time $t + \Delta t$.

The assumption that collisions only cause small changes in the particle's velocity implies that the probability $F(v, \Delta v)$ peaks around $\Delta v = 0$. Expanding F to second order in Δv will lead to the Fokker-Planck collision operator in one dimension. Including all three directions in velocity space, over which we count the indices i, j , the collision operator takes the following form

$$C\{f\} = \left(\frac{\partial f}{\partial t} \right)_c = \sum_{i,j} \frac{\partial}{\partial v_i} \left[- \frac{\langle \Delta v_i \rangle}{\Delta t} f + \frac{\partial}{\partial v_j} \left(\frac{\langle \Delta v_i \Delta v_j \rangle}{2\Delta t} f \right) \right], \quad (2.26)$$

which includes both first and second Δv -moments of the transition probability distribution F . The operator consists of contributions of two different kinds of terms: the advective terms that include first order derivatives, describing a dynamical frictional force, and the terms for diffusion in velocity space with second order derivatives. The latter is defined in terms of a tensor, $\langle \Delta v_i \Delta v_j \rangle / 2\Delta t$, reflecting the fact that diffusion is not necessarily isotropic.

Generally, any plasma consists of electrons and ions of various isotopes and charge states. Each distribution function f_a of species a is affected by collisions with all other species b (and with a itself) in an additive way, i.e. $C_a\{f_a\} = \sum_b C_{ab}\{f_a, f_b\}$. Thus, every contribution from all species b to the total collision operator $C_a\{f_a\}$ of species a need to be accounted for. In our context, a very important contribution is that of electron-electron collisions, which we write as $C_{ee}\{f_e, f_e\}$. It is common to consider the electron distribution to be dominated by a thermal bulk of electrons, a fact that can be used to simplify this collision operator. The electron distribution function is separated into a Maxwellian f_0 and a small non-Maxwellian perturbation $f_1 = f_e - f_0$, which is assumed to contain a much smaller number of particles than the thermal bulk. The collision operator is bilinear [14], and can therefore be expanded into

$$C_{ee}\{f_e, f_e\} = C_{ee}\{f_0, f_0\} + C_{ee}\{f_0, f_1\} + C_{ee}\{f_1, f_0\} + C_{ee}\{f_1, f_1\}. \quad (2.27)$$

The contribution of collisions between two Maxwellians with the same temperature and mean velocity is zero, and thus the first term in Eq. (2.27) vanishes. Since the non-Maxwellian population is assumed to be small, the contribution from self-collisions in f_1 is negligible.

Out of the two remaining terms in Eq. (2.27), $C_{ee}\{f_0, f_1\}$ and $C_{ee}\{f_1, f_0\}$, let us consider the latter, which is often called the *test particle term*. It accounts for the effects on the non-Maxwellian electrons from collisions with the much larger thermal population. The thermal electrons are described by the Maxwell-Boltzmann distribution in Eq. (2.19), taken to have a zero mean flow $\mathbf{u} = 0$. Therefore, f_0 is isotropic in momentum space, i.e. $f_0(\mathbf{v}) = f_0(v)$, where $v = |\mathbf{v}|$ is the magnitude of the velocity (not to be confused with the one-dimensional velocity v in Eq. (2.25)). As shown in Ref. [14], the test particle term of the non-relativistic Fokker-Planck collision operator is given by

$$C_{ee}\{f_1, f_0\} = \nu_D^{\text{ee}} \mathcal{L}\{f_1\} + \frac{1}{v^2} \frac{\partial}{\partial v} \left[v^3 \left(\frac{1}{2} \nu_s^{\text{ee}} f_1 + \frac{1}{2} \nu_{\parallel}^{\text{ee}} v \frac{\partial f_1}{\partial v} \right) \right], \quad (2.28)$$

The slowing-down frequency ν_s^{ee} describes a dynamical friction due to collisions $F_{\text{drag}} \sim m_e v_{\parallel} \nu_s^{\text{ee}}$, and the parallel velocity diffusion frequency $\nu_{\parallel}^{\text{ee}}$ suppresses any sharp gradients in the energy distribution of the non-Maxwellian electrons. The Lorentz—or pitch angle scattering—operator $\mathcal{L}$ is proportional to the angular part of the Laplacian ∇^2 and describes diffusion on a sphere on constant v in momentum space. Therefore, the deflection frequency ν_D^{ee} acts to isotropize in f_1 . A similar dynamic is also present in collisions between electrons and the much more massive ions (thus assumed to be stationary), for which $C_{ei}\{f_1, f_i\} = \hat{\nu}_{ei} \mathcal{L}\{f_1\}$ with the thermal collision frequency $\hat{\nu}_{ei}$, given in Eq. (2.4), for electron-ion collision. All combined, these collisional effects drive the distribution f towards thermodynamical equilibrium, i.e. a Maxwellian.

By defining the so-called *Chandrasekhar function* as follows

$$\mathcal{G}(x) := \frac{\text{erf}(x) - x \text{erf}'(x)}{2x^2}, \quad (2.29)$$

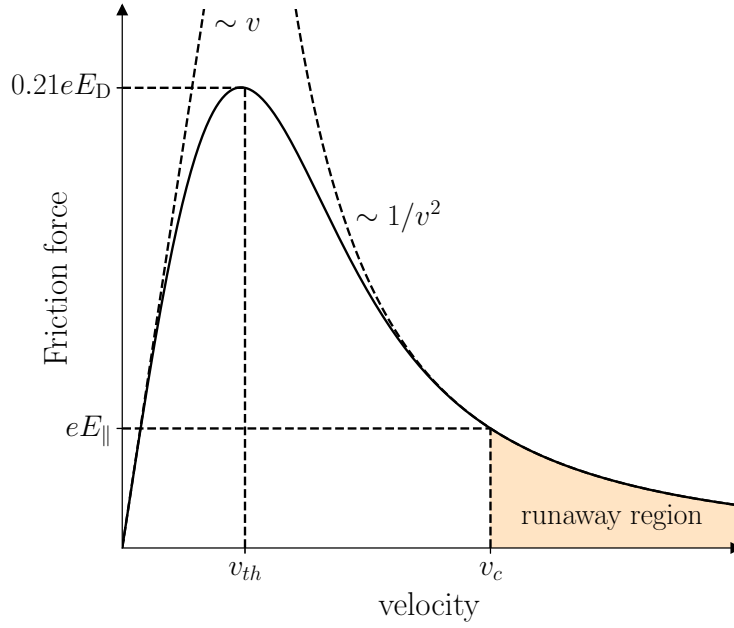


Figure 2.3: Illustration of the average friction force experienced by an electron moving in a plasma, as a function of its velocity. For electrons with $v < v_c$, the accelerating electric field $E_{\parallel}$ is counteracted by friction, while particles in the runaway region $v > v_c$ are accelerated towards relativistic speeds. For $E_{\parallel} > 0.21E_D$, where E_D is the Dreicer field, all electrons become runaways.

where $\text{erf}(x)$ denotes the error function¹, the collision frequencies in Eq. (2.28) are given by

$$\nu_D^{\text{ee}}(v) := \frac{\langle (\Delta v_{\perp}/v)^2 \rangle^{\text{ee}}}{2\Delta t} = \hat{\nu}_{\text{ee}} \frac{\text{erf}(x) - \mathcal{G}(x)}{x^3}, \quad (2.30a)$$

$$\nu_s^{\text{ee}}(v) := -\frac{\langle \Delta v_{\parallel}/v \rangle^{\text{ee}}}{\Delta t} = 4\hat{\nu}_{\text{ee}} \frac{\mathcal{G}(x)}{x}, \quad (2.30b)$$

$$\nu_{\parallel}^{\text{ee}}(v) := \frac{\langle (\Delta v_{\parallel}/v)^2 \rangle^{\text{ee}}}{\Delta t} = 2\hat{\nu}_{\text{ee}} \frac{\mathcal{G}(x)}{x^2}. \quad (2.30c)$$

where $x := v/v_{\text{th}}$, $\hat{\nu}_{\text{ee}}$ is the thermal collision frequency between two electron populations [14]. With this expression for the slowing-down frequency ν_s^{ee} in Eq. (2.30b), it follows that the dynamical friction is proportional to the Chandrasekhar function.

2.4.2 The runaway phenomenon

A particle moving through an ordinary fluid experiences a drag force that monotonically increases with its velocity, which is explained by the simple fact that it will come in contact with other particles at a higher rate when moving at a higher velocity, increasing the dynamical friction. Take for instance a falling object in a

¹The error function is defined as $\text{erf}(x) := \frac{2}{\sqrt{\pi}} \int_0^x ds e^{-s^2}$.

constant gravitational field; it will accelerate until the net force acting on it is zero, i.e. when the drag force becomes equal to the accelerating gravitational force, and then reach its maximal velocity.

The same is also true in plasmas for charged particles at low velocities, where the relative velocities between any two colliding particles is comparable to the thermal velocity. However, the dynamical friction decreases at superthermal velocities since $F_{\text{drag}} \propto \mathcal{G}(v/v_{\text{th}})$, as we saw in Sec. 2.4.1. Fig. 2.3 shows this force as a function of the particle velocity, from which this property is apparent; as $x \rightarrow \infty$, the Chandrasekhar function $\mathcal{G}(x) \rightarrow 1/2x^2$ asymptotically. Therefore, in the presence of an accelerating electric field $E_{\parallel}$, once the particle reaches a certain critical speed $v_c > v_{\text{th}}$, which is marked in Fig. 2.3, the force will not be able to counteract this acceleration and the particle runs away to extremely high energies. This is called the *runaway phenomenon*, and electrons with velocities $v > v_c$ is referred to as *runaway electrons* (REs). The first detailed account of this phenomenon was given by Dreicer [16, 17].

If the electric field is sufficiently strong for ordinary thermal electrons to become runaways, then the assumption that the non-Maxwellian part of the electron distribution is a small perturbation from the Maxwellian bulk population breaks down. Roughly, this occurs when the electric field becomes a sizable fraction of the so-called *Dreicer field* E_{D} , which is defined by

$$E_{\text{D}} := \frac{n_e e^3 \ln \Lambda}{4\pi \varepsilon_0^2 T_e}. \quad (2.31)$$

More specifically, it can be shown that all electrons will immediately become runaways for approximately $E_{\parallel} > 0.21 E_{\text{D}}$.

With the non-relativistic Fokker-Planck collision operator introduced in Sec. 2.4.1, the dynamical friction force approaches zero for particles with increasing velocities $v > v_{\text{th}}$. However, according to the special theory of relativity, the speed of light c is the upper limit for particles to travel at and one would thus expect that the friction does not drop all the way down to zero at such large energies. A relativistic treatment of runaways in the high energy limit shows that the friction force approaches the accelerating force corresponding to a critical electric field E_c , given by

$$E_c := \frac{T_e}{m_e c^2} E_{\text{D}} \quad (2.32)$$

assuming a fully ionized plasma [18]. For an electric field $E_c < E < 0.21 E_{\text{D}}$, only a fraction of all electrons will become runaways.

This theoretically predicted threshold for the accelerating electric field in Eq. (2.32) sometimes differs significantly from the measured “effective” critical electric field E_c^{eff} in various experimental setups [19]. The difference can be accredited to effects which enhance the drag force force experienced by the electrons, namely those of collisions with partially ionized atoms as well as synchrotron and bremsstrahlung losses [20–22].

Runaway generation mechanisms are usually divided into two categories: those which do not require an initial seed of runaways, called *primary generation*, and those which do, called *secondary generation*. The former includes the so-called *Dreicer mechanism*, which is a process of diffusive leaking of particles into the runaway

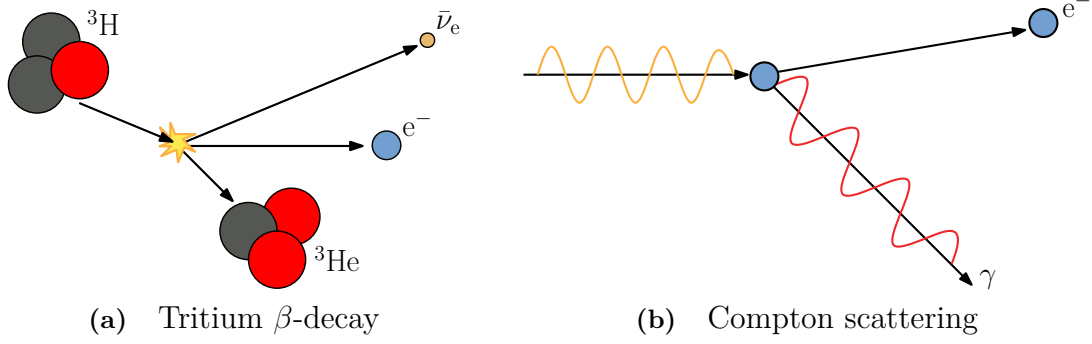


Figure 2.4: Relevant subatomic processes for RE seed generation during nuclear operations in reactor-scale tokamaks.

region. Recalling that collisions are driving the electron distribution towards thermal equilibrium, as mentioned in Sec. 2.4.1, thermal electrons will fill in the “gap” the REs leave behind and thus turn into runaways as well. The Dreicer mechanism is usually a major source of runaways in present day machines [23]. However, it is expected to play a relatively small role in ITER and even negligible when compared to other runaway generation mechanisms [24–26].

Another primary generation mechanism is the *hot-tail* mechanism, which occurs when a plasma is rapidly cooled down [27]. Since particles at high speeds are less affected by collisions than slower ones, they take more time to cool down and will therefore form a “tail” of superthermal (or “hot”) particles in the distribution function. Recalling from Sec. 2.1 that the conductivity of the plasma follows $\sigma \propto T^{3/2}$, a drop in temperature implies an increase in plasma resistivity, which in turn yields a temporary increase in the electric field before the electrons are able to respond. If the increase in $E_{\parallel}$ occurs over a sufficiently short time before the tail is able to lose its energy via collisions, the particles may be turned into runaways, given they move at speeds above the new decreased critical speed v_c . The hot-tail generation is thus strongly dependent on the rate of the temperature drop [28–30].

During nuclear operations, two seed generation mechanisms relevant for ITER are those of tritium β -decay and Compton scattering of electrons with photons from the neutron-activated wall, as illustrated in Fig. 2.4. When the plasma contains tritium, which has a half-life of $\tau_T = 4500 \pm 8$ days, electrons are produced as it decays into helium-3 via the decay process

$${}^3\text{H} \xrightarrow{\beta^-} {}^3\text{He} + e^- + \bar{\nu}_e \quad (18.6 \text{ keV}), \quad (2.33)$$

where $\bar{\nu}_e$ denotes an electron antineutrino. If a sufficient amount of the freed energy is released via the emitted β -electron, it turns into a runaway. Furthermore, the generation of runaways due to Compton scattering is caused by γ -rays emitted from the radioactive wall components facing the plasma, being activated by the barrage of neutrons released in the DT fusion reactions. These photons can interact with electrons in the plasma and scatter them over the runaway threshold.

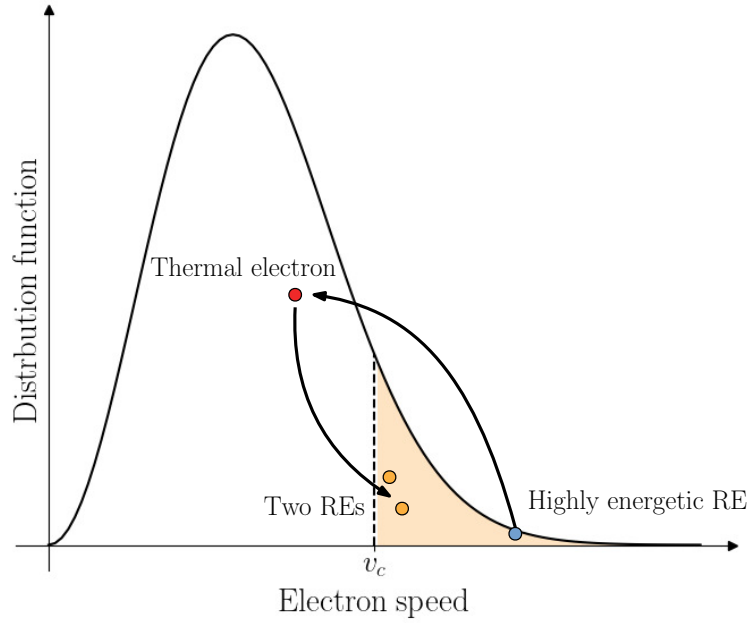


Figure 2.5: Illustration of a highly energetic RE, colliding with a much slower thermal electron such that both end up as REs, effectively multiplying its contribution to the RE density by a factor of two. The runaway region in the electron distribution is where the speed of an electron $v > v_c$.

The seed of runaways produced from these four mechanisms can subsequently be multiplied through knock-on collisions via the *avalanche* mechanism, as illustrated in Fig. 2.5. Such large-angle collisions are explicitly neglected in Fokker-Planck theory, briefly discussed in Eq. (2.26), which is derived under the assumption that the momentum transfer in each collision is small. When a highly energetic RE collides with a thermal electron it is possible for it to transfer a significant fraction of its energy, such that both it and the target electron become runaways. Even if it is slowed down to a speed just above v_c , the electric field quickly accelerates it back towards relativistic energies. The rate of which knock-on collisions occur is proportional to the amount of runaways already present in the plasma, resulting in an exponential growth of the highly energetic electrons [31]. Avalanche generation is expected to be responsible for the vast majority of runaways in ITER [24, 25].

2.5 The tokamak

So far the topics of discussion have been more or less general concepts within plasma physics. In this section we apply the theory and equations that have been introduced, in order to give a more detailed description of tokamak reactors. First, in Sec. 2.5.1, an overview of the tokamak design is provided. Section 2.5.2 then goes on to explain the different interactions that take place during a disruption event.

2.5.1 Fields and particle motion

Since charged particle trajectories follow magnetic field lines, the latter can be used to confine a plasma within an enclosed volume. In a tokamak, this is achieved by bending the field lines into the shape of a torus, using external magnetic coils. The resulting field is toroidally symmetric and can therefore be described by the *Grad-Shafranov equation* [32, 33]. In this description, the magnetic field lines follow surfaces of constant poloidal magnetic flux, ψ , which also corresponds to surfaces of constant magnetic pressure P (introduced in Eq. (2.24)). Furthermore, we assume that the gyro-radius of all species are smaller than the length scales of interest, such that $r_L \ll L_B$, implying that particles follow the magnetic field lines very closely. As such, a particle's position will be uniquely determined by the poloidal flux ψ , poloidal angle θ , and toroidal angle φ . A more comprehensive description of this geometry can be found in appendix A.

Using a purely toroidal magnetic field would however not be sufficient for magnetic confinement. The magnetic field is necessarily non-uniform, as it needs to be curved to fit in a finite volume, and its divergence-free nature implies then a spatial variation of the field strength. Consequently, there is a gradient pointing towards the axis of rotational symmetry, entailing a particle drift described by Eq. (2.12), in which ions and electrons will be displaced vertically in opposite directions. Furthermore, the resulting charge separation gives rise to an electric field, which introduces an additional drift mechanism according to Eq. (2.13), causing the particles to be transported radially outward from the center of the tokamak. In order to counteract these effects, the magnetic field lines are therefore twisted poloidally, such that they follow helical paths. The poloidal motion causes the ∇B effect to be averaged out, since the drift will be directed both towards- and away from the center of the plasma during a single revolution. In order to induce the poloidal magnetic field, a large current on the order of ~ 15 MA in ITER, as indicated in Fig. 1.1, is conducted through the plasma. It is in turn generated by passing a current, that is linearly varying in time, through a central solenoid, driving a current in the plasma playing the role of the secondary loop of a transformer. As long as the electron distribution is close to a Maxwellian, the current density follows *Ohm's law*, given by

$$\mathbf{j}_\Omega = \sigma \mathbf{E}, \quad (2.34)$$

where $\mathbf{j}_\Omega$ denotes the Ohmic current density.

An important observation is that while particle motion is tied to surfaces of constant poloidal flux, the field strength along magnetic field lines varies, decreasing at greater distances from the center of the tokamak. This quality gives rise to the *magnetic mirror* effect. Consider the kinetic energy, W_{kin} , of a particle moving along the twisted field lines with a parallel velocity $v_{\parallel}$. From the expression for the magnetic moment in Eq. (2.10), we may write this as

$$W_{\text{kin}} = \frac{mv_{\parallel}^2}{2} + \mu B, \quad (2.35)$$

where energy conservation requires that $W_{\text{kin}} = \text{const.}$ In addition, the adiabatic invariance of the magnetic moment similarly implies that $\mu = \text{const.}$ Consequently,

as the magnetic field strength increases, the parallel velocity must decrease. Assume that the lowest field strength along a particle's trajectory is given by B_0 , and denote the velocity at this point $v_0^2 = v_{\parallel,0}^2 + v_{\perp,0}^2$. If the maximum field strength along the orbit is $B_{\max}$, then from Eq. (2.35) it can be shown that the parallel velocity will be fully depleted if $v_0^2 B_0 = v_{\perp,0}^2 B_{\max}$. As such, for all particles with a velocity on the low field side of the flux surface satisfying

$$\frac{v_{\perp,0}^2}{v_{\parallel,0}^2 + v_{\perp,0}^2} \geq \frac{B_0}{B_{\max}}, \quad (2.36)$$

the parallel velocity $v_{\parallel}$ will change sign during the orbit, causing the particle to bounce back in what is called a *trapped orbit*. Since these particles will keep bouncing back and forth, they are "trapped" in the sense that they do not experience a net acceleration from any electric field parallel to the magnetic field lines, because the overall effect will cancel out over one bounce period. This holds under the assumption that the time scale of the acceleration caused by the electric field is significantly longer than that of the trapped orbit.

2.5.2 Disruptions

As described in Sec. 1.2, one of the largest obstacles for the tokamak concept concerns sudden loss of plasma confinement, which can lead to excessive heat loads on the plasma facing wall components and the generation of REs. The temporary loss of energy confinement is related to MHD instabilities and often emerge from excessive pressure gradients or current densities [34]. These critical conditions can for example be the result of a sudden influx of impurities, such as material being exfoliated from the wall (or injected purposefully). In these events, the flux surfaces are broken up, giving rise to intermediate structures in the form of *magnetic islands*, before eventually reconnecting in a chaotic manner. The process results in rapid particle- and heat transport that causes the plasma temperature to drop substantially. Furthermore, while the transport is less prominent at lower temperatures, if the temperature becomes sufficiently low, any impurities (as well as the fuel in some scenarios) may recombine and create an additional heat loss channel in the form of line radiation. The temperature dependence of radiated power for different neutrals and ion species is non-monotonic, but does in general exhibit a significant increase as the temperature decreases to the ~ 100 eV range. The positive feedback loop from this can lead to a *radiative collapse*, in which the remaining heat is very rapidly exhausted even further, yielding a temperature in the order of ~ 10 eV. The described phenomenon is often referred to as the *thermal quench* (TQ), and the heat that is transported during this phase poses a major issue for the wall integrity on its own.

In addition to potential heat losses, during the TQ there are mechanisms that complicate things even further. As was noted in Sec. 2.1, the conductivity of the plasma is proportional to the temperature as $\sigma \propto T^{3/2}$. The sudden temperature drop across three orders of magnitude is therefore accompanied by an even greater drop in conductivity. Considering Eq. (2.34), one could expect that the Ohmic current would decrease during the TQ as a result. However, the tokamak plasma

is an inductive system and any change in the plasma current is therefore counteracted through the subsequent generation of an electric field. The process can be inferred from Maxwell's equations. Specifically, consider taking the time derivative of Eq. (2.6d). Using Eq. (2.6c), the left-hand side yields

$$\nabla \times \frac{\partial \mathbf{B}}{\partial t} = -\nabla \times (\nabla \times \mathbf{E}) = \nabla^2 \mathbf{E}. \quad (2.37)$$

For the last equality we used the vector calculus identity $\nabla^2 \mathbf{E} = \nabla(\nabla \cdot \mathbf{E}) - \nabla \times (\nabla \times \mathbf{E})$, under the assumption that quasi-neutrality makes it such that, for the concerned length scales, $\nabla \cdot \mathbf{E} = 0$. Additionally we assume that the time scale of this interaction is significantly longer than that of electromagnetic waves, justifying the displacement current to be neglected. Including the right-hand side of the equation we then obtain

$$\nabla^2 \mathbf{E} = \mu_0 \frac{\partial \mathbf{j}}{\partial t}. \quad (2.38)$$

As a result, an electric field is induced that acts to preserve the magnitude of the Ohmic current during the TQ. Eventually the electric field will dissipate, along with the Ohmic current, throughout what is known as the *current quench* (CQ). The time scale of this dissipation is rather significant from a stability perspective, owing to the possibility of eddy and halo currents which can lead to magnetic forces being exerted on the tokamak device. It is therefore useful to define the *current quench time*, τ_{CQ} , according to

$$\tau_{\text{CQ}} := \frac{t(I_{\Omega} = 0.2I_{\text{p}}^{(t=0)}) - t(I_{\Omega} = 0.8I_{\text{p}}^{(t=0)})}{0.6}, \quad (2.39)$$

where I_{Ω} denotes the Ohmic current and $I_{\text{p}}^{(t=0)}$ is the total plasma current at the start of the disruption.

Finally, the induced electric field can accelerate electrons to relativistic velocities through the phenomenon described in Sec. 2.4.2, converting part of the Ohmic current into runaways. During the TQ and in the beginning of the CQ, it is therefore common that a seed of REs is produced by the primary runaway generation mechanisms, and the resulting current is then amplified on the account of the avalanche mechanism. Since the drag force experienced by the REs is very small, the runaway current will linger after the CQ, slowly decaying during what is known as the runaway plateau. The rate of the decay is determined by the effective critical field, $E_{\text{c}}^{\text{eff}}$, and the magnitude of the RE current [35].

Figure 2.6 shows the electron temperature T_{e} , induced electric field $E_{\parallel}$, as well as the Ohmic-, RE- and total plasma current I_{Ω} , I_{re} and I_{p} respectively, during the described phases of a disruption. It can be seen that I_{Ω} even increases slightly during the TQ², and then diminishes completely during the CQ. At the same time the induced field, $E_{\parallel}$, gives rise to a growing runaway current, that eventually starts to decay very slowly during the runaway plateau.

²This so called " I_{p} spike" is outside of the scope of this thesis, but it can be modeled in DREAM; the interested reader is referred to Ref. [37]

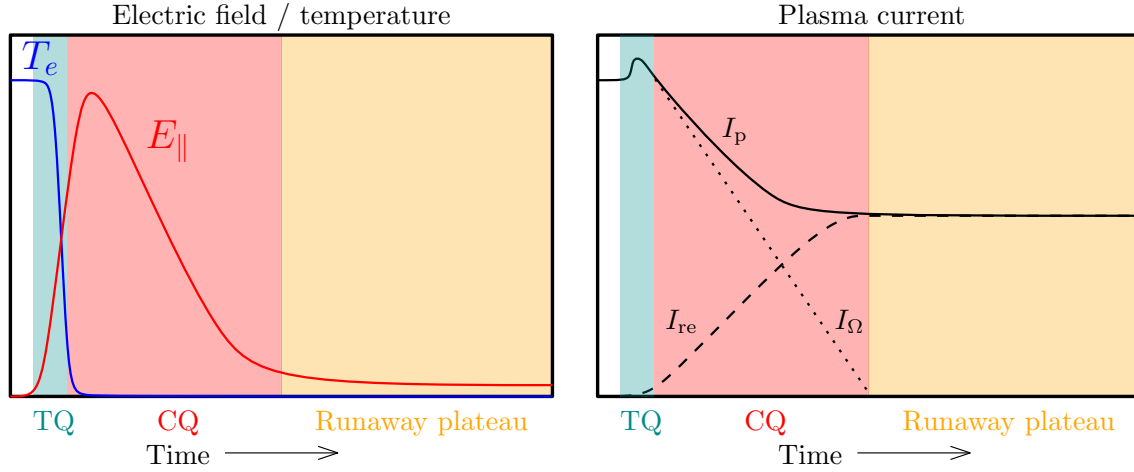


Figure 2.6: Illustration of the electric field and temperature (left), as well as the different plasma currents (right) during a disruption. During the TQ, the current density is radially redistributed leading to a temporary increase in I_{Ω} . I_{Ω} later dissipates during the CQ, and in part being converted into a runaway current that lingers afterwards in the runaway plateau. Source: [36].

As previously stated, a disruption scenario introduces issues related to heat transport, CQ time, and runaway currents. A definition of τ_{CQ} was given in Eq. (2.39), but in order to make the other properties more explicit we define

$$I_{RE}^{\max} := \max_t I_{RE}(t) \quad (2.40)$$

$$\eta_{\text{cond}} := \frac{W_{\text{cond}}}{W_{\text{th}}^{(t=0)}}. \quad (2.41)$$

Eq. (2.40) simply defines the maximal runaway current, $I_{RE}^{\max}$, achieved during the disruption. In Eq. (2.41), η_{cond} represents the transported fraction of the initial thermal energy, with W_{cond} denoting the amount of heat that has been transported out through the plasma edge and $W_{\text{th}}^{(t=0)}$ being the thermal energy of the electron population at $t = 0$. For a disruption scenario to be tolerable in an ITER-like tokamak, these quantities must be limited such that [38, 7]:

$$I_{RE}^{\max} < 150 \text{ kA}, \quad \eta_{\text{cond}} < 10 \%, \quad 50 \text{ ms} < \tau_{CQ} < 150 \text{ ms}. \quad (2.42)$$

Finally, an additional constraint that is used in this work regards the Ohmic current. If the final Ohmic current, $I_{\Omega}^{\text{final}}$, towards the end of the disruption is sufficiently large, it can lead to issues related to the mentioned halo current and runaway conversion. Since a complete termination of tokamak operation is desired in the event of a disruption, $I_{\Omega}^{\text{final}}$ should preferably be negligible.

To satisfy the given conditions, most disruptions cannot be left unmitigated. For this purpose, mitigation techniques are employed in an attempt to moderate the mechanical and thermal stress exerted on the device. These methods are however still in development and it remains unclear whether it will be possible to adequately suppress disruption events in reactor-scale tokamaks. Next section will provide a brief overview of the topic.

2.5.2.1 Mitigation techniques

The most widely used mitigation technique, that also has been chosen as the primary method for dealing with disruptions in the ITER tokamak [38], is *massive material injection* (MMI). In this scheme one would inject new material into the tokamak at the onset of a disruption. The material would then be able to radiate away much of the thermal energy through line radiation. This is preferable since it can reduce localised heat loads, while also affecting the overall temperature evolution which in turn impacts the Ohmic CQ time. Finally, since the injected material quickly ionizes, the total number of free electrons also increases. This amplifies the drag force experienced by the electrons and as a result suppresses runaway generation. These qualities make MMI a promising concept, as it would allow for some control over all of the mentioned complications that occur during a disruption. With that being said, as we will see shortly MMI is associated with its own set of difficulties.

Due to their radiative properties, and for not being chemically reactive, noble gases such as neon are primary candidates for the injection. There are however reasons for having other types of material as well. As the plasma cools down the ions begin to recombine. It has been shown that bound electrons contribute less to the drag force than free electrons due to atomic screening, while their contribution to the avalanche growth rate is similar to that of the free electrons [23]. This implies that the injected material could potentially increase runaway generation as well. In order to counteract this, injection schemes can also include lighter elements, such as hydrogen isotopes. These do not radiate energy as efficiently (as they tend to get fully ionized) but they do not recombine in the same extent as noble gases and can therefore be a more successful runaway generation suppressant.

Another mitigation technique is the use of external magnetic perturbations. By inducing particle and heat transport in a controlled manner, it could be possible to eject REs from the plasma before the avalanche mechanism can produce a substantial runaway current, while also reducing the localized heat loads. This appears to be viable in smaller devices as in SPARC [39]. On the other hand there is a problem with core penetration on ITER size devices. It is feared that the REs generated in the center of the plasma will not come into contact with the perturbed regions on the edge, meaning they will not be transported [40]. This is especially problematic since runaway generation tend to be more prominent near the center. The method is therefore not as commonly used in practice as MMI.

In this chapter we have covered some of the underlying theory regarding plasmas in general, as well as in the context of tokamak reactors. The next chapter will instead focus on how these concepts are implemented in simulations in order to adequately model a disruption scenario.

3

Disruption model

In this work we employ the recently developed *Disruption Runaway Electron Analysis Model* (DREAM) tool [41] when simulating a range of disruption mitigation scenarios for an ITER-like plasma. It is a finite-volume fluid-kinetic framework used to model RE dynamics during tokamak disruptions, which allows treating the electrons at different degrees of approximation. As we need to be able to run a large number of disruption simulations during a limited amount of time, when performing any of the black-box optimizations described in Ch. 4, we limit our modelling to the simplest, fluid treatment of the plasma. More specifically, this includes modelling the thermal bulk of “cold” electrons and the small runaway population as two separate fluid species. The former is characterized by a density n_{cold} , a temperature T_{cold} as well as an Ohmic current density j_{Ω} and the REs are described by their density n_{RE} . It is assumed that the runaways move with the speed of light parallel to the magnetic field, hence their associated current density is $j_{\text{RE}} = ec n_{\text{RE}}$.

In fluid mode, the code resolves one spatial dimension, namely a flux surface label r as radial coordinate, defined as the half width of a flux surface in the midplane (where $z = 0$ in Cartesian coordinates). Thus, the quantities evaluated by DREAM are flux surface averages. We consider a near-Maxwellian plasma in force balance between the pressure and magnetic forces, as expressed in Eq. (2.24). It is assumed that the trajectories of charged particles follow magnetic field lines *exactly*, moving on surfaces of constant poloidal flux ψ ; the radial coordinate r being a flux surface label implies that it is held fixed along a particle orbit. DREAM implements a parameterization of the magnetic field $\mathbf{B}$, fully described by its major radius, elongation, Shafranov shift, triangularity and a reference poloidal flux profile. In this work, we consider an ITER-like plasma, which is described more in detail in appendix A.

This chapter outlines the system of equations used in this work to simulate various tokamak disruption scenarios. First, the evolution of the total current density is described in Sec. 3.1, followed by the dynamics of the background plasma in Sec. 3.2. The latter includes the thermal electron population as well as the density $n_i^{(j)}$ of each charge state j of every ion species i present within the plasma. Furthermore, as described in Sec. 3.3, the time evolution of the REs is determined by the transport due to magnetic perturbations and a sum of five source terms, each of which is a model for describing the generation of runaways as the result of a few different physical phenomena. Finally, a brief overview on how DREAM self-consistently solves this system of equations is presented in Sec. 3.4.

3.1 Current density evolution

The evolution of the current density is calculated by evolving the poloidal flux ψ via the equation

$$\frac{\partial \psi}{\partial t} = 2\pi \frac{\langle \mathbf{E} \cdot \mathbf{B} \rangle}{\langle \mathbf{B} \cdot \nabla \varphi \rangle} =: -V_{\text{loop}}, \quad (3.1)$$

which is derived from Faraday's law of induction. Here, $\langle \cdot \rangle$ denotes the flux surface average, defined in appendix A, and φ is the toroidal angle. The loop voltage, V_{loop} in Eq. (3.1), is related to j_Ω through Ohm's law

$$\frac{j_\Omega}{B} = \sigma \frac{\langle \mathbf{E} \cdot \mathbf{B} \rangle}{\langle B^2 \rangle}, \quad (3.2)$$

where $\sigma = \sigma(n_{\text{cold}}, T_{\text{cold}}, Z_{\text{eff}})$ denotes the parallel electric conductivity and is calculated using the Sauter-Redl model, which accounts for neoclassical effects across all collisionality regimes [42]. The *effective charge number* Z_{eff} is a property of the background plasma and is defined in Eq. (3.7).

Ampère's law relates the total parallel current density $j_{\text{tot}} := j_{\text{RE}} + j_\Omega$ to the poloidal flux through

$$2\pi\mu_0 \langle \mathbf{B} \cdot \nabla \varphi \rangle \frac{j_{\text{tot}}}{B} = \frac{1}{V'} \frac{\partial}{\partial r} \left[V' \left\langle \frac{|\nabla r|^2}{R^2} \right\rangle \frac{\partial \psi}{\partial r} \right], \quad (3.3)$$

where the Jacobian $V' = V'(r)$ is the total area of any flux surface at r , defined in appendix A. Multiplying this area with the infinitesimal distance dr , yields the incremental volume between two surfaces at r and $r + dr$, respectively.

At $t = 0$ no runaways are present and the initial total current density profile is prescribed as $j_{\text{tot}}(r, 0) = j_\Omega(r, 0) \propto [1 - (r/a)^2]^{0.41}$, with a total plasma current of $I_p = 15$ MA, which sets the initial condition for Eq. (3.1). It is defined as the total toroidal current enclosed within the plasma $I_p := I(a)$, where

$$I(r) := \frac{1}{2\pi} \int_0^r V' dr' \frac{j_{\text{tot}}}{B} \langle \mathbf{B} \cdot \varphi \rangle, \quad (3.4)$$

and $a = 2$ m is the minor radius of the plasma.

A boundary condition for Eq. (3.3) is obtained by assuming that the plasma is surrounded by a perfectly conducting wall at $r = b$, leaving $\psi(b) = 0$. Here, we set $b = 2.15$ m. The poloidal flux at the plasma edge is coupled to the wall via an edge-wall mutual flux inductance M_{ew} as $\psi(a) = \psi(b) - M_{\text{ew}} I_{\text{tot}}$.

3.2 Background plasma dynamics

We consider the total electron density $n_{\text{tot}} = n_{\text{cold}} + n_{\text{RE}}$, including the larger thermal electron density n_{cold} and a smaller runaway population $n_{\text{RE}} \ll n_{\text{cold}}$. At the start of a disruption simulation it is assumed that there are no REs present inside the plasma. Being in thermal equilibrium, equipartition of energy yields the average

energy density w_{cold} of the thermal population,

$$w_{\text{cold}} := \frac{3}{2} n_{\text{cold}} T_{\text{cold}}. \quad (3.5)$$

To ensure that quasi-neutrality is satisfied, n_{cold} is evolved in time such that the overall charge density is zero, i.e.

$$n_{\text{cold}} = \sum_i \sum_{j=0}^{Z_i} Z_{0j} n_i^{(j)} - n_{\text{RE}}. \quad (3.6)$$

Here Z_i denotes the ion atomic number and $Z_{0j} := j$ is the charge number. Subtracted from the sum of ionic charges is the density of the runaway population, n_{RE} , which is evolved as described in Sec. 3.3. Having introduced the ion densities, the effective charge number is defined as

$$Z_{\text{eff}} := \sum_i \sum_{j=0}^{Z_i} Z_{0j}^2 n_i^{(j)} / \sum_i \sum_{j=0}^{Z_i} Z_{0j} n_i^{(j)}, \quad (3.7)$$

which is a measure of the average charge number of all the background ions.

Each ion charge state density is evolved via ionization and recombination processes, neglecting charge-exchange terms with neutral hydrogenic species. More specifically, it will follow the evolution of the ions given by

$$\frac{\partial n_i^{(j)}}{\partial t} = n_{\text{cold}} \left(\mathcal{I}_i^{(j-1)} n_i^{(j-1)} - (\mathcal{I}_i^{(j)} + \mathcal{R}_i^{(j)}) n_i^{(j)} + \mathcal{R}_i^{(j+1)} n_i^{(j+1)} \right), \quad (3.8)$$

where $\mathcal{I}_i^{(j)}$ and $\mathcal{R}_i^{(j)}$ are the corresponding ionization and recombination rates respectively, both of which are dependent on temperature. The species of ions considered within this work are those of the hydrogen isotopes, deuterium and tritium, as well as the ~ 10 times heavier noble gas neon. For the neon ions, the rates $\mathcal{I}_i^{(j)}$ and $\mathcal{R}_i^{(j)}$ are extracted from the ADAS database [43] and similarly from the AMJUEL database [44] for the hydrogen species. The latter accounts for the opacity to Lyman radiation, which becomes significant at the high hydrogen densities considered here. Prior to $t = 0$, the plasma fuel is assumed to consist of an even mixture of deuterium and tritium, both of which are initially radially uniform and fully ionized. At $t = 0$, a material injection of neutral deuterium and neon takes place with specified initial radial profiles. For simplicity, we disregard the thermal energy of the ions.

There are two possible methods for evolving the temperature of the cold electrons in fluid simulations with DREAM, both of which are used in this work. The first is to directly prescribe a radial temperature profile as a function of time. While this provides good control over the simulation, there are certain interactions and processes that will not be included as a result. A more complete approach is to evolve the temperature self-consistently. Through the relation given in Eq. (3.5), the change in temperature can be inferred from the energy density equation, presented in Eq. (2.23), with the assumed form

$$\frac{\partial w_{\text{cold}}}{\partial t} = j_{\Omega} E_{\parallel} + j_{\text{RE}} E_{\text{c}} - n_{\text{cold}} \sum_i \sum_{j=0}^{Z_i} n_i^{(j)} L_i^{(j)} + \frac{1}{V'} \frac{\partial}{\partial r} \left[\frac{3n_{\text{cold}}}{2} V' D_W \frac{\partial T_{\text{cold}}}{\partial r} \right]. \quad (3.9)$$

The first term on the right hand side describes Ohmic heating of the system, caused by the induced electric field, where $E_{\parallel} := \langle \mathbf{E} \cdot \mathbf{B} \rangle / B$ is the component of the electric field parallel to the local magnetic field. The second term represents the heating of the cold electrons from collisions with REs, with E_c denoting the critical electric field, defined in Eq. (2.32). The third term is a sink/source term introduced to describe energy change resulting from atomic processes, with $L_i^{(j)}$ denoting the corresponding rate of energy loss due to ionization, recombination, line radiation and brehmsstrahlung. As with the ionization and recombination rates, $L_i^{(j)}$ has a temperature dependence and all relevant data regarding this quantity are extracted from the ADAS and AMJUEL databases.

The fourth term in Eq. (3.9) represents a diffusive heat transport. The diffusion coefficient, D_W , needs to be prescribed, which, as mentioned in Sec. 2.3.2, serves as a form of fluid closure. In this work we use the *Rechester-Rosenbluth diffusion model*

$$D = \pi R_0 |v_{\parallel}| \left(\frac{\delta B}{B} \right)^2, \quad (3.10)$$

to model both the electron heat and the runaway electron particle transport. The equation describes radial transport caused by turbulent perturbations of the magnetic field [45]. Here, R_0 is the major radius, $\delta B / B$ is the normalized amplitude of the magnetic perturbations, and $v_{\parallel}$ is the velocity of the particle along the magnetic field lines. Note that the magnetic perturbation strength can vary both in space and time. Here, $\delta B / B \neq 0$ only in Sec. 5.2 and 5.3, and the effect of transport of both heat and REs is otherwise neglected. Finally, in order to compute the heat diffusion coefficient, D_W , the heat flux moment of Eq. (3.10) is taken with the distribution function assumed to be a Maxwellian.

Since tokamak disruptions are characterized by a rapid temperature decrease, accompanied by a rapid decrease in conductivity, it is essential that this process, i.e. the TQ, is modeled in an appropriate manner. The initial temperature drop is dominated by transport losses, which we cannot model fully self-consistently. For this phase we therefore need to either prescribe the temperature altogether or prescribe the magnetic perturbation strength. Given an initial profile, $T_{\text{init}}(r)$, and a final profile, $T_{\text{final}}(r)$, one possibility is to prescribe an exponential decay according to

$$T_{\text{cold}}(r, t) = T_{\text{final}} + [T_{\text{init}}(r) - T_{\text{final}}] \exp(-t/t_0), \quad (3.11)$$

with the decay time $t_0 = 1$ ms. The temperature profiles used here are $T_{\text{init}} = T_0[1 - 0.99(r/a)^2]$, with $T_0 = 20$ keV, and the radially uniform $T_{\text{final}} = 10$ eV. These values are inspired by previous ITER studies [25]. Prescribing the evolution based on Eq. (3.11), then switching to a self-consistent configuration after the plasma has lost most of its thermal energy at $t = 8$ ms, is a viable approach to inducing a disruption scenario in a DREAM simulation. This approach works well under the condition that there is enough injected material and the temperature is sufficiently low at the onset of the self-consistent evolution to yield a radiative collapse, ensuring that the temperature does not begin to increase right away.

Alternatively, it is possible to set the initial temperature $T_{\text{cold}}(r, 0) = T_{\text{init}}(r)$

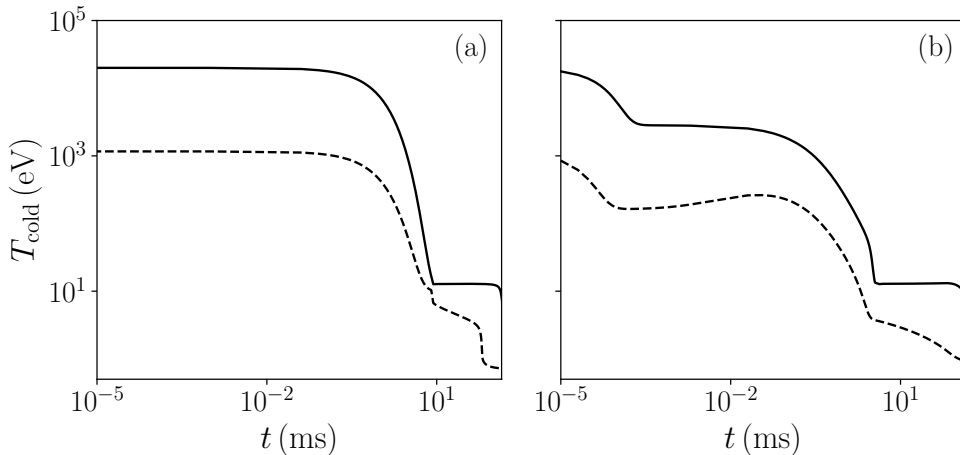


Figure 3.1: Temperature evolution at the center (solid) as well as the edge (dashed) of the plasma during a disruption simulation, where deuterium and neon were uniformly injected at the start of the simulation with densities $n_{\text{D}} = 6 \cdot 10^{20} \text{ m}^{-3}$, $n_{\text{Ne}} = 6 \cdot 10^{17} \text{ m}^{-3}$. The temperature during the TQ is either prescribed (a), based on Eq. (3.11), or it is evolved self-consistently (b), using Eq. (3.9) and (3.10) assuming $\delta B/B = 0.3\%$. After the completion of the TQ (~ 8 ms in (a) and ~ 4 ms in (b)), a self-consistent temperature evolution with no heat transport is used in both cases. Note the log-log scale.

and have a self-consistent evolution of T_{cold} throughout the entire simulation. This does however introduce additional constraints for the simulation parameters, if an adequate TQ is to be achieved. The transport efficiency depends on the magnetic perturbation strength, $\delta B/B$, provided to the diffusion coefficient in Eq. (3.10). These perturbations are a rather complex phenomenon, the study of which is outside the scope of this project. A simple model is therefore assumed for this quantity, in which the perturbations are uniform in space. The values used will be based on previous observations that the perturbations are significantly higher during the TQ, where they in certain scenarios can exceed 1%, than during the remainder of the disruption [46]. As such, it is assumed that the perturbation strength is constant during the TQ until the temperature has dropped below $T_{\text{cold}} < 20 \text{ eV}$, and then zero during the CQ phase of the simulations, meaning there is no heat transport.

The two methods for obtaining a TQ are compared in Fig. 3.1, where the temperature evolution at the center and edge ($r = 0$ and $r = a$) of the plasma are shown for a characteristic disruption simulation, in which deuterium and neon were uniformly injected in the beginning with densities $n_{\text{D}} = 6 \cdot 10^{20} \text{ m}^{-3}$, $n_{\text{Ne}} = 6 \cdot 10^{17} \text{ m}^{-3}$. The simulation in Fig. 3.1a used a prescribed temperature according to Eq. (3.11). The TQ was in this case completed after ~ 8 ms, after which the simulation is switched to a self-consistent evolution. An analogous simulation, but with a self-consistent temperature configuration, is shown in Fig. 3.1b. Here it was assumed that $\delta B/B = 0.3\%$ during the TQ, which finished after ~ 4 ms.

It is interesting to note the different behavior exhibited during the self-consistent TQ in panel b, which can be tied to the different processes presented in Eq. (3.9). Ini-

tially, there is a very rapid ($\sim 0.1 \mu\text{s}$) temperature drop resulting from the ionization of the injected deuterium and neon. After the material has been fully ionized, the dominating process for the temperature evolution is the diffusive transport, which occurs on a much longer time scale. As previously stated, the transport was turned off at $\sim 4 \text{ ms}$, and the subsequent decrease in plasma temperature at the edge of the plasma compared to the center is likely a result of the Ohmic heating power being stronger in the core. The same is true at the conclusion of the prescribed temperature evolution in panel a.

In addition to offering a more complete physical description of the temperature evolution, the self-consistent approach also allows for direct computation of certain quantities of interest. One example that is especially relevant in the context of disruptions is the amount of energy that is transported out of the plasma, W_{cond} . Assuming that all heat transported across the last closed flux surface ($r = a$) is lost to the tokamak wall, it is given by

$$W_{\text{cond}} := \int dt \left(\frac{3n_{\text{cold}}}{2} V' D_W \frac{\partial T_{\text{cold}}}{\partial r} \right) \Big|_{r=a}. \quad (3.12)$$

This will be utilized later in Ch. 5, when quantifying the performance of disruption simulations through Eq. (2.41).

3.3 Runaway electron generation

The time evolution of REs is determined by radial transport due to magnetic perturbations, as well as a number of runaway generation sources that feed particles directly into the runaway population. The flux surface averaged RE density is evolved according to

$$\begin{aligned} \frac{\partial n_{\text{RE}}}{\partial t} = & \left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{Dreicer}} + \left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{hot-tail}} + \left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{tritium}} \\ & + \left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{Compton}} + \Gamma_{\text{ava}} n_{\text{RE}} + \frac{1}{V'} \frac{\partial}{\partial r} \left[V' D_{\text{RE}} \frac{\partial n_{\text{RE}}}{\partial r} \right]. \end{aligned} \quad (3.13)$$

The Dreicer mechanism is accounted for in the term $(\partial n_{\text{RE}}/\partial t)^{\text{Dreicer}}$, which is calculated by a neural network trained on data from kinetic simulations over a range of plasma parameters and impurity concentrations [47].

For the second term in Eq. (3.13), DREAM implements a fluid model for hot-tail generation based on Ref. [48]. Assuming a high effective charge Z_{eff} , the hot-tail runaway rate $(\partial n_{\text{RE}}/\partial t)^{\text{hot-tail}}$ is given by

$$\left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{hot-tail}} = -4\pi p_c^2 \frac{\partial p_c}{\partial t} f_0(r, p_c). \quad (3.14)$$

Here, the isotropic electron distribution function f_0 is based on the non-relativistic slowing-down distribution derived in Ref. [49], with the initial condition of being a Maxwellian distribution with a density $n_{\text{cold}}(r)$ and temperature $T_{\text{cold}}(r)$ at $t = 0$. The critical momentum p_c is obtained by considering the angle-averaged Boltzmann

equation (Eq. (2.17)) in the Lorentz limit of strong pitch-angle scattering. This yields the following isotropic source term

$$S(p) = \frac{1}{3} \left(\frac{E_{\parallel}}{E_c} \right)^2 \frac{1}{1 + Z_{\text{eff}}} \frac{p^3}{\gamma} \frac{\partial f_0}{\partial p} + \frac{\gamma^2}{p^2} f_0, \quad (3.15)$$

where $\gamma := (1 - v^2/c^2)^{-1/2}$ denotes the Lorentz factor and E_c is the critical electric field for runaway generation, as defined in Eq. (2.32). A positive $S(p')$ implies a net flow of electrons into the momentum sphere $p < p'$. Since f_0 is monotonically decreasing, the first source term in Eq. (3.15) is always negative and acts as an acceleration of electrons uniformly in all directions. This term dominates for large p . The second term is always positive and acts as a friction force that slows down the electrons, which dominates for small p [48]. The critical momentum is defined as $p_c : S(p_c) = 0$, i.e. the momentum that balances the two terms in Eq. (3.15) with one another, from which p_c is solved for numerically.

The third runaway generation in Eq. (3.13) is modeled as in Refs. [24, 26] and accounts for the REs generated when tritium decays into helium-3. As mentioned in Sec. 2.4.2, only the produced electrons with an energy higher than the critical runaway energy, W_{crit} , will contribute to the generation of runaways in $(\partial n_{\text{RE}}/\partial t)^{\text{tritium}}$. This critical energy is given by $W_{\text{crit}} := m_e c^2 (\sqrt{p_*^2 + 1} - 1)$, in terms of the critical momentum (normalized to $m_e c^2$) for runaway acceleration, which is obtained from the implicit expression

$$p_* = \frac{\sqrt[4]{\bar{\nu}_s(p_*) \bar{\nu}_D(p_*)}}{\sqrt{E_{\parallel}/E_c}}. \quad (3.16)$$

In this equation $\bar{\nu}_s$ and $\bar{\nu}_D$ are relativistic slowing-down and deflection collision frequencies normalized to $\nu_c = eE_c/m_e c$, accounting for the energy dependence of the Coulomb logarithm and the effect of partial screening in collisions with partially ionized impurities. These are sums of the corresponding contributions from electron-electron and electron-ion collisions, as well as an additional term in the slowing-down frequency, which accounts for the radiation losses due to brehmsstrahlung [22]. In a fully ionized plasma, these normalized collision frequencies $\bar{\nu}_s \rightarrow 1$ and $\bar{\nu}_D \rightarrow 1 + Z_{\text{eff}}$.

For a plasma with a tritium density of n_T , the generation of runaways due to tritium decay is thus modeled as

$$\left(\frac{\partial n_{\text{RE}}}{\partial t} \right)^{\text{tritium}} = \ln(2) \frac{n_T}{\tau_T} F_{\beta}(W_{\text{crit}}). \quad (3.17)$$

Here, $F_{\beta}(W_{\text{crit}})$ is the fraction of electrons produced in this process with energies above the critical runaway energy W_{crit} . This fraction is determined by an integral over the energy spectrum of the beta-electrons and can be approximated by treating a beta-electron as a free particle, i.e. neglecting its Coulomb interaction with the produced helium-3 nucleus, yielding the expression $F_{\beta}(W_{\text{crit}}) \approx 1 - (35/8)x^{3/2} + (21/4)x^{5/2} - (15/8)x^{7/2}$, with $x := W_{\text{crit}}/Q$ and $Q = 18.6\text{keV}$ being the total energy released from the decay process and thus the maximum kinetic energy possible for the beta-electron [26].

The fourth growth rate $(\partial n_{\text{RE}}/\partial t)^{\text{Compton}}$ models the Compton scattering of electrons into the runaway region in momentum space, according to Ref. [24]. Given as an integral over the photon energy ε_γ , it is determined by

$$\left(\frac{\partial n_{\text{RE}}}{\partial t}\right)^{\text{Compton}} = n_{\text{tot}} \int_0^\infty d\varepsilon_\gamma \Gamma_\gamma(\varepsilon_\gamma) \sigma(\varepsilon_\gamma), \quad (3.18)$$

where Γ_γ is the estimated photon energy spectrum and σ the total Compton scattering cross section [24]. The energy of the photons are in general much larger than the ionization energy of the ions and atoms in the plasma, which is the reason for this growth rate being proportional to the total density of electrons n_{tot} rather than that of the free electrons n_{free} .

The fifth runaway source term accounts for the avalanche mechanism. The avalanche growth rate is modeled as in Ref. [23], given by

$$\Gamma_{\text{ava}} = \frac{e}{m_e c \ln \Lambda_c} \frac{n_{\text{tot}}}{n_{\text{cold}}} \frac{E_{\parallel} - E_c^{\text{eff}}}{\sqrt{4\bar{\nu}_s(p_*) + \bar{\nu}_s(p_*)\bar{\nu}_D(p_*)}}, \quad (3.19)$$

where E_c^{eff} is the effective critical electric field [22], which takes into account screening effects and losses due to brehmsstrahlung and synchrotron radiation. In order to obtain a well-behaved formula for when the accelerating electric field is lower than E_c^{eff} , the critical momentum in Eq. (3.16) is calculated by replacing $p_*(E_{\parallel})$ with $p_*(E_c^{\text{eff}})$ for $E_{\parallel} < E_c^{\text{eff}}$. This is used to approximately describe the runaway decay for $E_{\parallel} \approx E_c^{\text{eff}}$ in the absence of losses due to magnetic perturbations.

Finally, the last term in Eq. (3.13) describes the transport of REs due to magnetic perturbations. Whenever $\delta B/B$ is non-zero, we prescribe such transport by using the Rechester-Rosenbluth diffusion coefficient D_{RE} , calculated using Eq. (3.10) with the runaways moving along the field lines at the speed of light $v_{\parallel} = c$.

3.4 Numerical implementation

In this chapter we have introduced a number of equations used when simulating tokamak disruptions in DREAM. These need to be solved numerically, which is not a trivial task considering that it involves the computation of several integro-differential expressions. For this purpose, DREAM uses a *finite volume discretization*. In this framework, the phase space parameters are discretized into separate cells and the equation system is expressed in terms of the cell averaged quantities, as well as their values at the cell faces. These values make up two overlapping grids, which we denote the *flux grid* and *distribution grid* respectively. In the fluid model used in this work, quantities are only resolved on one-dimensional cells, corresponding to a radial grid of size N_r . Considering some quantity $X(r)$, the cell average values, associated with the cell centers, are denoted X_i with integer indices $i = 1, \dots, N_r$. Face values located on the flux grid are instead denoted using half integer indices $i \pm 1/2 = 1/2, \dots, N_r + 1/2$. When required, these are evaluated through interpolation on the distribution grid. This discretization is illustrated in Fig. 3.2.

The flux grid points are defined using a cell width, Δr_i , from which one also

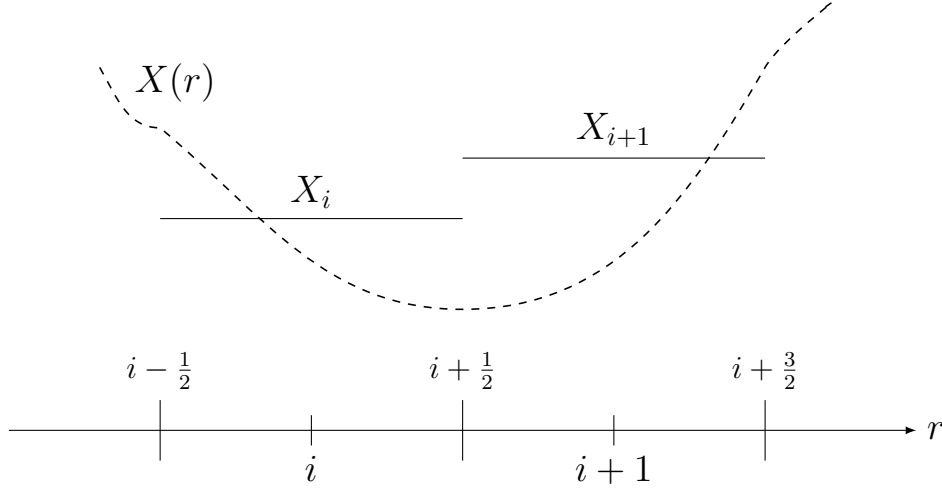


Figure 3.2: Illustration of the grid cells used when discretizing the radial coordinate system. The distribution grid, denoted by integer indices, contains the cell averaged value of the function $X(r)$, while the flux grid, denoted by half integer indices, is located at cell faces, where the fluxes across grid cells are evaluated.

derives the distribution grid points according to

$$r_{i+1/2} = r_{i-1/2} + \Delta r_i, \quad (3.20)$$

$$r_i = \frac{r_{i+1/2} + r_{i-1/2}}{2}. \quad (3.21)$$

Note that the width may vary for different cells. In this scheme, a central difference approximation is used for the spatial derivatives, while the implicit Euler method is used for time derivatives [41]. Letting X_i^ℓ denote the central value in cell i at timestep ℓ , these can be written as

$$\left(\frac{\partial X}{\partial r}\right)_i^\ell \approx \frac{X_{i+1/2}^\ell - X_{i-1/2}^\ell}{\Delta r_i}, \quad (3.22)$$

$$\left(\frac{\partial X}{\partial t}\right)_i^{\ell+1} \approx \frac{X_i^{\ell+1} - X_i^\ell}{\Delta t^\ell}, \quad (3.23)$$

where Δt^ℓ is the size of timestep ℓ . Due to the fact that fluxes are computed using the values on the flux grid, as shown in Eq. (3.22), the flux into a grid cell is numerically equal to the flux out of the adjacent cells. This assures a machine-precision conservation of the integrated value for X [50].

When running simulations, all of the essential relations presented in Secs. 3.1-3.3 are formulated in terms of this discretization. Considering that separate equations are obtained for each point along the distribution grid, this results in a rather large system of non-linear equations. Let k denote the total number of equations and $\mathbf{x}^\ell \in \mathbb{R}^k$ be the corresponding array of unknown quantities at timestep ℓ . The equations can then be written in the compact form

$$\mathbf{F}^\ell(\mathbf{x}^{\ell+1}) = 0, \quad (3.24)$$

where $\mathbf{F}^\ell : \mathbb{R}^k \rightarrow \mathbb{R}^k$ is a vector function that varies for each time step to account for the implicit Euler scheme used. Eq. (3.24) is a root finding problem, which is solved iteratively using *Newton's method* [51]. The iteration scheme is given by

$$\mathbf{x}_{n+1}^{\ell+1} = \mathbf{x}_n^{\ell+1} - J_{\mathbf{F}}^{-1}(\mathbf{x}_n^{\ell+1}) \mathbf{F}^\ell(\mathbf{x}_n^{\ell+1}), \quad (3.25)$$

where n denotes the current iteration step (not to be confused with the phase space index), and $J_{\mathbf{F}}^{-1}$ is the inverse Jacobian matrix with respect of $\mathbf{F}^\ell$ with respect to $\mathbf{x}$. The Jacobian can be derived analytically, but a number of approximations are also used during its construction for the sake of sparsity and simplicity [41]. Newton's method is assumed to have converged when

$$\|x_{n+1} - x_n\|_2 \leq \epsilon_x^{\text{abs}} + \epsilon_x^{\text{rel}} \|x_{n+1}\|_2, \quad \forall x \in \mathbf{x}^{\ell+1}. \quad (3.26)$$

Here $\|\cdot\|_2$ is the 2-norm, while ϵ_x^{abs} and ϵ_x^{rel} denote the absolute- and relative tolerances for the unknown quantity x respectively. In this work, the default tolerances were set to $\epsilon_0^{\text{abs}} = 0$ and $\epsilon_0^{\text{rel}} = 10^{-2}$, while the RE density used the tolerances $\epsilon_{\text{nRE}}^{\text{abs}} = 10^5 \text{ m}^{-3}$ and $\epsilon_{\text{nRE}}^{\text{rel}} = 2 \cdot 10^{-6}$. Furthermore, the radial grid was resolved with $\Delta r = 0.1 \text{ m}$, whereas the time resolution varied such that $\Delta t \sim 10^{-10} - 10^{-5} \text{ s}$. The latter was done to assure that the initial ionization of the injected material, as well as the small time scale interactions during the TQ, were properly resolved.

Finally it should be noted that while Newton's method is used to solve a set of non-linear equations, each iteration step of Eq. (3.25) also involves solving a set of linear equations. This is done through LU factorization, implemented by the parallel direct solver PARDISO, which is part of Intel's Math Kernel Library (MKL)¹.

This concludes our discussion about disruption simulations, as implemented by the DREAM code. The topic of next chapter is optimization and the techniques that we intend to apply to this framework.

¹Link to the MKL library: <https://www.intel.com/content/www/us/en/developer/tools/oneapi/onemkl.html>

4

Black-box optimization

When faced with the problem of optimizing a function $f : \mathcal{X} \rightarrow \mathbb{R}_+$, with $\mathbb{R}_+$ being the positive real numbers, on the subspace $\mathcal{P} \subseteq \mathcal{X}$, there are numerous techniques that one can choose from with varying degrees of efficiency and informational gain. In this project, certain properties of a tokamak plasma disruption represent the subject of optimization. As described in Ch. 3, the disruption is represented by a complex numerical model described by a system of strongly nonlinear integro-differential equations with a large number of variables. There are of course countless ways of setting up such a simulation, but what they have in common is that they take a vector $\mathbf{x}$ of parameters as input, and compute a set of relevant output quantities from which the function $y = f(\mathbf{x})$ is obtained. Referred to as the *objective function*, f is designed to depend on these simulation output quantities such that small values of the function correspond to more favorable outcomes, while larger values are undesirable. The two objective functions used within this work are described in Sec. 5.1 and 5.2.

The task at hand is encapsulated by the more general minimization problem of finding the optimum $\mathbf{x}^*$, defined by

$$\mathbf{x}^* := \operatorname{argmin}_{\mathbf{x} \in \mathcal{P}} f(\mathbf{x}). \quad (4.1)$$

where $\mathcal{P} := \{\mathbf{x} \in \mathcal{X} \mid a_i \leq x_i \leq b_i, i = 1, \dots, d\}$ is a hyper-rectangle on which the optimization is performed, and d is the number of input parameters considered, i.e. the dimensionality of $\mathcal{P}$. Its lower and upper bounds need to be appropriately chosen in order for the global optimum to be contained within $\mathcal{P}$.

Furthermore, the function f is assumed to be expensive to evaluate, in the sense that there is but a limited number of simulations that are practical/feasible to perform. In this work, this is a justified assumption as a single evaluation takes a substantial amount of time to compute, usually ~ 5 minutes, depending on the input settings. This circumstance limits the practicality of using many of the conventional methods of optimization, especially those relying on the ability to determine derivatives of f , e.g. gradient descent (requiring gradient) and Newton's method (requiring both gradient and Hessian). Even though the problems considered within this thesis are deterministic, more general stochastic problems can be formulated with the optimization framework¹ developed during this project.

One possible recourse is to use what is sometimes referred to as a "black-box" approach. This means that there is no assumption made about any special structure of the function, e.g. concavity or linearity, that would make the task of optimization

¹Link to Github repository: <https://github.com/peterhalldestam/exjobb.git>

much easier using other techniques that leverages such structures to improve the efficiency of the method [52]. This approach is well motivated by the mere complexity of the underlying system of equations considered within this work. In this scheme the function is viewed simply as something that takes some input and produces a corresponding output.

In this chapter we consider two different approaches for finding the optimum defined in Eq. (4.1), namely the *Powell's conjugate direction method* and *Bayesian optimization*, which are described in Sec. 4.1 and 4.2, respectively. The first approach is to search for an optimum with as few function evaluations as possible by performing a number of one-dimensional optimization subroutines in a progression of directions. The second is based on Bayesian statistics, where knowledge about the objective function is gathered by exploring $\mathcal{P}$ globally and sequentially updating a probability distribution of f , which is used when deciding the next input parameter to sample.

4.1 Powell's conjugate direction method

While optimization in multiple dimensions entails greater complexity and computational cost than a one-dimensional scenario, the fundamental procedure can be made quite similar by dividing the process into a series of one-dimensional minimizations. Starting at a point $\mathbf{x} = \mathbf{P} \in \mathcal{P}$ and given some directional vector $\mathbf{u} \in \mathbb{R}^d$, one can instead express the single variable function $g : \mathbb{R} \rightarrow \mathbb{R}_+$ along a single line within $\mathcal{P}$ as

$$g(\lambda) := f(\mathbf{P} + \lambda \mathbf{u}), \quad \lambda \in \mathbb{R}. \quad (4.2)$$

A one-dimensional optimization method can then be applied to Eq. (4.2) to find the point, $\mathbf{P}_{\min} = \mathbf{P} + \lambda_{\min} \mathbf{u}$, along the specified line for which f is minimized. Given instead a full set of basis vectors $\{\mathbf{u}_i\}_{i=1}^d$ that span $\mathbb{R}^d$, one can imagine sequencing a number of line minimizations in different directions in such a way that the routine eventually converges toward an optimum in $\mathbb{R}^d$ -space. This is the main idea behind the so-called *direction set methods*. The discussion of these, and further concepts presented in this section, follow the reasoning given in Ref. [51].

A very straightforward approach to the direction set technique would be to assume a basis of orthogonal unit vectors $\{\mathbf{e}_i\}_{i=1}^d$, and then cycle through them, optimizing for each direction and continuously moving the starting point $\mathbf{P} \rightarrow \mathbf{P} + \lambda_{\min} \mathbf{e}_i$, until the function stops decreasing. This approach does in fact work well for many functions, but there are exceptions. If the level surfaces of the function happens to create narrow “corridors” at an angle with respect to the chosen basis, then the routine will have trouble navigating the parameter space resulting in many small steps, thus making the algorithm inefficient. This complication is illustrated in Fig. 4.1. In general, these cumbersome surfaces are mathematically characterized by the function having a second derivative that is greatly varying depending on the direction.

The mentioned problem could be circumvented if appropriate directions for optimization were to be found. The question is then what directions qualify as “appro-

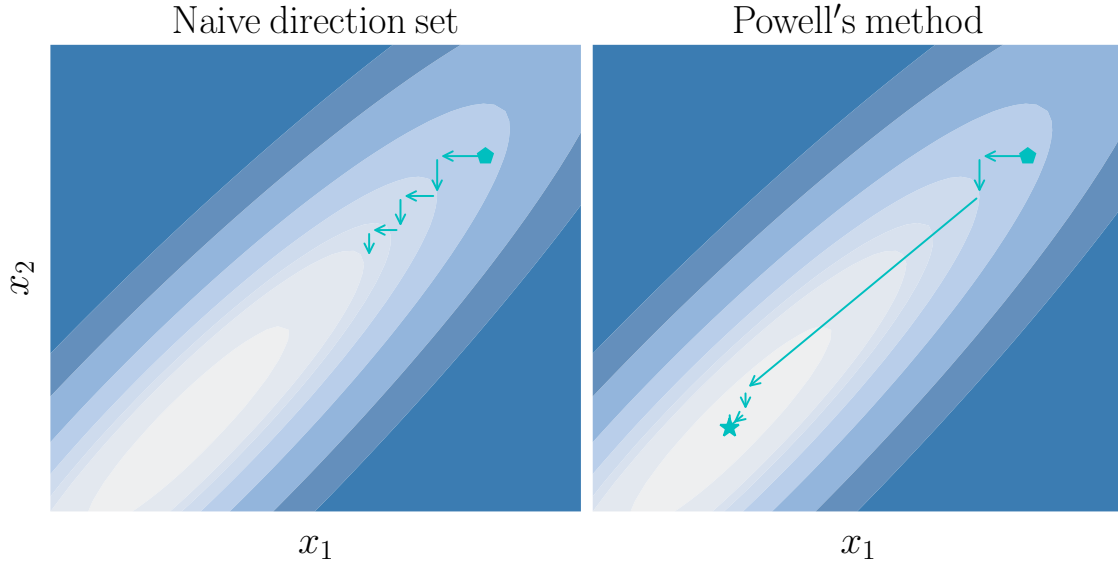


Figure 4.1: Illustration of six iterations with a naive direction set method and Powell's method in a region where the function surface creates a small “corridor”. The starting point is denoted by a pentagon and each line minimization is represented by an arrow. It can be seen that the naive approach struggles to make progress, while Powell's method manages to find a new advantageous direction and converges towards a single point denoted by a star.

appropriate”. One idea is to find a set of directions that do not interfere with each other, such that minimization along one direction is not counteracted by the subsequent minimization in another direction. These are known as *conjugate directions* and their existence can be explicitly derived. Consider the function f from before at some starting point $\mathbf{P}$. Using Taylor expansion one finds that to the second order, the function can be approximated as

$$f(\mathbf{x}) \approx f(\mathbf{P}) + \mathbf{b}^\top \mathbf{x} + \frac{1}{2} \mathbf{x}^\top H \mathbf{x}, \quad (4.3)$$

where

$$b_i := \left. \frac{\partial f}{\partial x_i} \right|_{\mathbf{P}} \quad H_{ij} := \left. \frac{\partial^2 f}{\partial x_i \partial x_j} \right|_{\mathbf{P}} \quad (4.4)$$

are the gradient and Hessian matrix respectively. The gradient of Eq. (4.3) is then

$$\nabla f = \mathbf{b} + H\mathbf{x}. \quad (4.5)$$

Consider now two vectors, $\mathbf{u}$ and $\mathbf{v}$. These are said to be conjugate if the gradient in one direction is invariant under translation along the other. In other words, when minimizing along $\mathbf{u}$ the change in gradient, $\delta(\nabla f)$, should be orthogonal to $\mathbf{v}$. From Eq. (4.5) one finds that $\delta(\nabla f) = H\delta\mathbf{x} = H\mathbf{u}$, assuming a translation where $\delta\mathbf{x} = \mathbf{u}$. The condition is then met if and only if

$$\mathbf{v}^\top \delta(\nabla f) = \mathbf{v}^\top H\mathbf{u} = 0. \quad (4.6)$$

The previous derivation confirms that conjugate directions do exist, what remains is to find them by generating a set of vectors that satisfy Eq. (4.6). This can be done using *Powell's method* outlined in Algorithm 1 [53]. In this schematic description the one-dimensional optimization method is not specified, but is instead represented by an arbitrary `linemin` function. The specific line minimization technique used in this work is discussed in Sec. 4.1.2.

Algorithm 1 Powell's method

```

Assume a starting point  $\mathbf{P}_0 \in \mathbb{R}^d$  and an initial basis set  $\{\mathbf{u}_i\}_{i=1}^d = \{\mathbf{e}_i\}_{i=1}^d$  that
spans  $\mathbb{R}^d$ .
while not converged do
  for  $i = 1, \dots, d$  do
    Find  $\lambda_{\min} = \text{linemin}(f(\mathbf{P}_{i-1} + \lambda \mathbf{u}_i))$ 
    Let  $\mathbf{P}_i = \mathbf{P}_{i-1} + \lambda_{\min} \mathbf{u}_i$ 
  end for
  Remove  $\mathbf{u}_1$  and offset the other basis vectors such that  $\mathbf{u}_i \rightarrow \mathbf{u}_{i-1}$ .
  Let  $\mathbf{u}_d = \mathbf{P}_d - \mathbf{P}_0$ .
  Find  $\lambda_{\min} = \text{linemin}(f(\mathbf{P}_d + \lambda \mathbf{u}_d))$ .
  Set  $\mathbf{P}_0 = \mathbf{P}_d + \lambda_{\min} \mathbf{u}_d$ .
end while
return  $\mathbf{P}_0$  and  $f(\mathbf{P}_0)$ .

```

Powell realized that conjugate vectors could be produced using the effective distance traveled after having cycled through a set basis vectors and minimized along each direction. He was able to show that if the function is of quadratic form, such as the one in Eq. (4.3), then one could obtain a new conjugate direction after each cycle by creating the vector $\mathbf{u} = \mathbf{P}_d - \mathbf{P}_0$, where d is the dimension of the optimization space. It is interesting to note how this procedure scales with the number of free parameters. From Algorithm 1 we gather that creating a new basis vector involves $d + 1$ line minimizations. In theory, the exact optimum for a quadratic function would be obtained once a complete set of conjugate directions is achieved. Since this requires d new basis vectors, Powell's method would then have a quadratic complexity $\mathcal{O}(d^2)$. In practice of course, most functions do not have this quadratic form. Most functions are however parabolic in the vicinity of a minimum. Because of this Powell's method will eventually start to *converge quadratically*. This means that near the minimum the number of significant digits approximately doubles for each iteration. While the quadratic convergence makes Powell's method very desirable, there is one issue that needs to be remedied before the routine can be used. This is the problem of linear dependence between the basis vectors, which will be the focus of the next section.

4.1.1 The problem of linear dependence

Powell's method offers good performance and is easily applicable to a black-box function. However, there is a major robustness issue in the routine described in Algorithm 1. When removing $\mathbf{u}_1$ in favor of the new direction there is a possibility

of the basis vectors becoming linearly dependent. Once this happens, parts of the initially considered parameter space $\mathcal{P} \subseteq \mathbb{R}^d$ will become inaccessible to the optimization routine and as such, the obtained minimum will only be that of a subspace $\mathcal{P}^* \subset \mathcal{P}$. An unpredictable reduction of the parameter space during optimization is not something that can be tolerated, but fortunately there are techniques that can be used to avoid this.

One simple way of making sure that the routine does not get stuck with linearly dependent vectors is to reset the basis to orthogonal unit vectors $\{\mathbf{e}_i\}_{i=1}^d$ every time d new directions have been found. This preserves the quadratic convergence property and is therefore a good alternative when the function is almost quadratic or high precision near the minimum is desired. It should however be noted that the property is not essential for all applications, and experience shows that this is the case for the disruption simulations in this work. Firstly, there is no evidence implying that the objective function behaves quadratically with respect to the different input parameters. Secondly, high accuracy is not imperative to the optimization, since the goal of this project is to find viable parameter regions during disruptions, while allowing for certain margins of error. Furthermore the objective function, while being physically motivated, is an *ad hoc* construction which implies that having a large number of significant digits in this application may not be more informative in a general context.

For the reasons described we concede the notion of quadratic convergence, and opt for a more heuristic method that aims to find a few good directions, rather than a full set of conjugate vectors. The idea is to only update the basis under certain conditions and in those scenarios remove the direction in which the function saw the largest decrease. The reason to discard this seemingly good direction instead of the arbitrary $\mathbf{u}_1$ is that it is likely a major component in the new vector and its removal will therefore reduce the chance of linear dependence occurring. What remains is to determine when the addition of a new direction is appropriate. Let $f_0 := f(\mathbf{P}_0)$ be the function value at the starting point and $f_d := f(\mathbf{P}_d)$ the value at the final point after a complete cycle of the basis. We also extrapolate an additional point $f_E := f(2\mathbf{P}_d - \mathbf{P}_0)$ corresponding to one full step in the proposed direction, and let Δf denote the largest decrease obtained in any one direction during the cycle. The new direction is included if:

1. $f_E < f_0$
2. $2(f_0 - 2f_d + f_E)(f_0 - f_d - \Delta f)^2 < (f_0 - f_E)^2 \Delta f$

The first condition makes sure that there is something to be gained from traveling in this direction and it is not “used up”. The second condition is more subtle. If the inequality is not fulfilled it can be argued that either the average decrease was not primarily a result of the decrease in any single direction, or that the second derivative of the function along the new direction is noticeably high and the current position appears to be close to its minimum.

This concludes the modification to prevent linear dependence and our discussion about Powell’s method in general. What has yet to be mentioned is how the one-dimensional minimization, so far only denoted `linemin`, is conducted. This is the topic of the next section.

4.1.2 Golden section search and Brent's method

Regardless of all the intricacies associated with direction set methods, what remains after the function has been defined along a line, as in Eq. (4.2), is a typical one-dimensional minimization problem. Here, there are again numerous techniques to choose from, with the caveat that they should not require computation of the derivative. In this project we have opted to use *Brent's method* which provides a flexible routine that performs well even for “ill-behaving” functions, while also converging quickly near the minimum [51]. It can be seen as an extension of the *golden section search* (GSS) method which is therefore the first target of discussion.

In order to find a minimum using GSS, it must first have been bracketed. The triplet (a, b, c) , where $a < b < c$, is said to bracket a minimum if it holds that $f(b) < f(a)$ and $f(b) < f(c)$. Exact details about how this bracket is obtained will be presented in the next section, but for now let us assume that it is given. Similarly to the bisection method used to find roots of functions, the aim is to reduce the bracketed interval by evaluating the function at some new point $x \in (a, b) \cup (b, c)$. For example, if the evaluated point is in the interval (a, b) one would update the bracket as follows:

- If $f(x) > f(b)$ then a is discarded and the new bracket will be (x, b, c) .
- If $f(x) < f(b)$ then it is the best minimum found so far and the interval can be reduced by discarding c , obtaining (a, x, b) .

This simple process can be repeated until the bracket becomes small enough to satisfy some convergence condition, after which the middle point of the remaining interval is returned as the location of the minimum.

There is still the question of how the new point x is chosen. It is preferable to make the safest choice possible such that the average number of required iterations, even in the worst case scenario, will be minimized. Let $\beta \in (0, 1)$ denote the fraction of the distance between a and c where b is located. Similarly, assume that x is located an additional fraction $\xi \in (0, 1)$ beyond b . More explicitly we have

$$\beta := \frac{b - a}{c - a} \tag{4.7}$$

$$\xi := \frac{x - b}{c - a}. \tag{4.8}$$

When evaluating $f(x)$, the two possible outcomes imply that the length of the new bracket will either be a fraction $\beta + \xi$, or $1 - \beta$ of the previous one. In order to minimize the length of the worst outcome these two should be equal, i.e.

$$\xi = 1 - 2\beta. \tag{4.9}$$

From Eq. (4.9) it follows that x will lie in the larger of the two intervals (a, b) and (b, c) , since $\xi < 0$ for $\beta > 1/2$. Something interesting happens when this selection is repeated. Assuming that the previous bracket was created using the same process, then β would be the same as the fraction of the distance between b and c where x is located. That is

$$\frac{\xi}{1 - \beta} = \beta. \tag{4.10}$$

Eq. (4.9) together with Eq. (4.10) yields a quadratic equation with the solution $\beta \approx 0.38197$. This means that the optimal bracket has its intermediate point b a fractional distance $\Gamma = 0.38197$ from one end and $1 - \Gamma = 0.61803$ from the other. These are the fractions of the *golden section* and they relate to the golden ratio, φ , as $2 - \varphi$ and $\varphi - 1$ respectively; hence the name *golden section search*.

GSS is a robust method which is suitable for minimizing practically any function. There are however techniques that converge faster, assuming that the function is sufficiently "well behaved". For example, if the function has close to a parabolic shape, then given the same bracket as before, one could imagine fitting the three points to some second degree polynomial and selecting the minimum of that polynomial as the next point to evaluate. It can be shown that this minimum is given by

$$x = b - \frac{1}{2} \frac{(b-a)^2[f(b) - f(c)] - (b-c)^2[f(b) - f(a)]}{(b-a)[f(b) - f(c)] - (b-c)[f(b) - f(a)]}. \quad (4.11)$$

This is known as *inverse parabolic interpolation*. If the studied function is in fact a second degree polynomial then this method will instantly find the minimum, and even if this is not the case it is still likely to provide a good guess. Recall that when sufficiently close to a minimum, most functions are parabolic. In a worst case scenario such a routine can however perform very poorly, constantly choosing its next step based on assuming a parabola which simply is not there. This is where Brent's method, outlined in Algorithm 2, offers a good middle ground.

Algorithm 2 Brent's method

```

Provide a triplet of points  $(a, b, c)$  that brackets the minimum.
while not converged do
    Determine a trial point  $\tilde{x}$  by attempting a parabolic fit according to Eq. (4.11).
    if  $\tilde{x}$  satisfies conditions in Eq. (4.12) and (4.13) then
        Accept the trial point, let  $x = \tilde{x}$ .
    else
        Take a golden section step, letting  $x = b + \Gamma(c - b)$  if  $(b, c)$  is the larger
        segment, or  $x = b - \Gamma(b - a)$  if not.
    end if
    Evaluate  $f(x)$  and update bracket accordingly.
end while
return  $b$  and  $f(b)$ .

```

What makes Brent's method so efficient is the fact that it is a combination of GSS and inverse parabolic interpolation, switching between the two depending on how the function is behaving. For each iteration a trial parabolic fit is done using the current bracket and a new point is suggested according to Eq. (4.11). The point is accepted if it satisfies the two following conditions:

$$x \in (a, c) \quad (4.12)$$

$$|x - b| < \frac{|\delta_2|}{2}. \quad (4.13)$$

Here δ_2 denotes the length of the step taken two iterations prior to the current one. Equation (4.12) is rather straightforward and it states that the new point should lie within the bracketed interval. The second constraint in Eq. (4.13) formulates that the attempted step should be at most half of δ_2 , which is done to confirm that the routine is converging and not stuck in some inefficient loop caused by a non-cooperative function. The reason for comparing the step length with that of two iterations ago and not simply the previous one is very pragmatic, as it has been found that allowing for one bad step can improve performance if it is compensated for by the following step. If these conditions are not met, then Brent's method will simply use a GSS step as detailed previously. One question remains before our description of the one-dimensional optimization is complete: how to bracket a minimum.

4.1.2.1 Bracketing a minimum

As with most concepts discussed in this chapter, there are several ways to bracket a minimum. There is however not much convention when it comes to this topic and we have therefore chosen a relatively simple method for our application. The idea is to progress “downhill”, increasing the size of each subsequent step, until the function begins to increase.

Consider an initial guess of two points (x, y) , where $x < y$. To begin, one evaluates the function at the respective points in order to see in which direction it is decreasing. For the sake of this illustration, let us assume that $f(x) > f(y)$. A new point z is then evaluated according to

$$z = y + \gamma L, \tag{4.14}$$

where $L := |x - y|$ is the length of the initial segment and γ is some hyperparameter governing the distance traveled. If $f(z) > f(y)$, then the process is complete and the bracket is given by $(a, b, c) = (x, y, z)$. If however $f(z) < f(y)$, then the segment is shifted such that $x \rightarrow y$ and $y \rightarrow z$, while the step size is doubled meaning that $\gamma \rightarrow 2\gamma$. The procedure is then repeated until the former inequality is true. This yields an exponential growth of γ , which is beneficial for the problem we consider, since in disruption simulations the function domain stretches over several orders of magnitude.

The initial guess could in practice be almost any two numbers, even randomly generated ones, but in this adaptation of the method they have been slightly tailored towards the studied parameter space. If we assume that the current line function is characterized by the point $\mathbf{P}$ and directional vector $\mathbf{u}$, then the line will intersect the boundary of the parameter space at the two points $\mathbf{P} + \ell^+ \mathbf{u}$ and $\mathbf{P} + \ell^- \mathbf{u}$, where $\ell^- < 0 < \ell^+$. It was found that the initial segment $(0, \ell^+/10)$ was a practical choice as it assures that the line minimization will in general consider a large part of the parameter space, while keeping the interval reasonably constrained with respect to the specified boundaries.

This concludes our discussion about one-dimensional optimization as well as Powell's method. There are however several reasons as to why a comparison with other techniques would be of interest. While efficiency is a primary concern of optimization methods, there are more aspects that should be taken into consideration. Most important is the fact that both Powell's method and Brent's method only finds a local minima, which can be a severe flaw depending on the shape of the studied function. Another aspect is the information gained during the procedure. Even if the function is minimized to a satisfactory extent, Powell's method will not tell us anything about the overall behavior of the function. Such knowledge is not only useful in the context of black-box optimization, but can be essential from a physics point of view. For these reasons we have chosen to use an additional optimization scheme known as *Bayesian optimization*.

4.2 Bayesian optimization

A practical feature of Bayesian statistics is the inference of the distribution of probabilities $p(H|\mathcal{D})$ for some hypothesis H given the dataset $\mathcal{D}$. This is guaranteed by *Bayes' theorem*, which allows for estimating this distribution in closed form, based on said data and any prior assumptions made about the hypothesis. Bayes' theorem can be derived from the definition of the conditional probability of some event A given another event B

$$p(A|B) := \frac{p(A \cap B)}{p(B)}, \quad (4.15)$$

where $p(A \cap B)$ is the probability for both events to occur (the intersection of the two). An analogous can be made for the reversed ordered $p(B|A)$ and since the intersection is commutative, this implies the following property of conditional probabilities

$$p(A|B)p(B) = p(B|A)p(A). \quad (4.16)$$

Solving for $p(A|B)$ and replacing A and B with the hypothesis $\mathcal{H}$ and dataset $\mathcal{D}$ respectively gives

$$p(\mathcal{H}|\mathcal{D}) = \frac{p(\mathcal{D}|\mathcal{H})p(\mathcal{H})}{p(\mathcal{D})}, \quad (4.17)$$

which is Bayes' theorem. Here, $p(\mathcal{H}|\mathcal{D})$ is known as the *posterior* probability distribution, expressing our knowledge about the hypothesis after observing each data point. On the right hand side of Eq. (4.17), the probability of observing $\mathcal{D}$ given the hypothesis is true is represented by $p(\mathcal{D}|\mathcal{H})$, and is called the *likelihood*. This describes the compatibility of the new data with this hypothesis. Furthermore, $p(\mathcal{H})$ is called the *prior*, which is our prior knowledge about the hypothesis. Finally, $p(\mathcal{D})$ in the denominator is called the *marginal likelihood* or *evidence*, and can simply be regarded as a normalization factor. If our hypothesis represents some value of the continuous variable $x \in \mathcal{X}$, where $\mathcal{X}$ is an exclusive and exhaustive set of outcomes,

this factor is set such that the posterior satisfies the normalization condition

$$\int_{\mathcal{X}} dx \, p(x|\mathcal{D}) = 1, \quad (4.18)$$

which is a fundamental property of probability distributions.

Provided a dataset $\mathcal{D}_n := (X_n, Y_n)$ consisting of n sample points in the sequence $X_n := \{\mathbf{x}_i\}_{i=1}^n$, and the corresponding function evaluations are $Y_n := \{f(\mathbf{x}_i)\}_{i=1}^n$, the use of Bayesian inference allows not only for approximating the objective function f , but also to infer the uncertainty of this approximation [54]. In linear regression of a parametric model, one could parameterize $f(\mathbf{x}) = \sum_i w_i \phi_i(\mathbf{x})$, given a set of known basis functions ϕ_i , $i = 1, \dots, M$ (for scalar inputs, this could for instance be a set of polynomials $\{1, x, x^2, \dots, x^{M-1}\}$). With this approach, it is then over the weights w_i that the inference is performed. The prior placed on f is obtained implicitly from any individual priors on the weights. In this work however, a non-parametric technique from Bayesian machine learning is used, which is described in Sec. 4.2.1.

Using Bayes' theorem in Eq. (4.17), we may express the inference of the function f from the training data in $\mathcal{D}_n$ as

$$p(f(\mathbf{x})|\mathcal{D}_n) = \frac{p(Y_n|f(\mathbf{x}), X_n)p(f(\mathbf{x}))}{p(\mathcal{D}_n)}. \quad (4.19)$$

This posterior distribution is sequentially updated by sampling more and more points until the algorithm is terminated. By carefully choosing where to sample next based on previous observations, as described in Sec. 4.2.2, one can balance the search in the vicinity of the currently evaluated optimum, with the exploration of regions with high uncertainty [55]. The latter makes it less likely that the algorithm gets stuck in local optima, which is a clear advantage over certain iterative methods, such as Powell's method described in Sec. 4.1. On the other hand, a larger number of evaluations is often required as the exploration needs to cover a greater part of parameter space.

This statistical approach also allows for dealing with certain aspects of f that are not of interest, known as *nuisance parameters*. It is thus a powerful tool for analyzing simulations with noisy data. If necessary, it is possible to include an independent and identically distributed Gaussian noise $\varepsilon \sim \mathcal{N}(0, \sigma_\varepsilon^2)$ with a zero mean and variance σ_ε^2 when evaluating $y = f(\mathbf{x}) + \varepsilon$. The kind of black-box functions considered within this thesis are noise-free and deterministic, i.e. each identically prepared simulation yields exactly the same outcome, we simply assume $\sigma_\varepsilon^2 = 0$.

Furthermore, as the *Gaussian process* (GP) can be instructed to explore regions in the vicinity of the optimum, it can therefore obtain a sufficiently accurate approximation of f in this region. This information can be used in order to get a picture of how sensitive the optimum is to small perturbations in $\mathbf{x}$ away from $\mathbf{x}^*$. This feature is used when analyzing the engineering robustness of the results in Ch. 5.

Summarized below in Algorithm 3 is the pseudo-code for the Bayesian optimization used within the works of this thesis [52]. It performs N function evaluations and then returns the point $\mathbf{x}^*$ with the smallest evaluated objective function. Two important and intricate elements of Algorithm 3 is the determination of the posterior distribution $p(f(\mathbf{x})|\mathcal{D}_n)$ and the acquisition function for intelligently selecting

the next point to sample. Section 4.2.1 describes how Bayesian inference is used to update the posterior, gathered from points selected as described in Sec. 4.2.2.

Algorithm 3 Bayesian optimization

Assume n_0 starting points $\{\mathbf{x}_i\}_{i=1}^{n_0}$.
 Evaluate initial dataset $\mathcal{D}_{n_0} = \{(\mathbf{x}_i, y_i)\}_{i=1}^{n_0}$.
 Place GP prior on f
 Let $n = n_0$
while $n \leq N$ **do**
 Update posterior distribution $p(f(\mathbf{x})|\mathcal{D}_n)$ using all available data in $\mathcal{D}_n$.
 Let $\mathbf{x}_{n+1}$ be the maximizer of the acquisition function EI_n on $\mathcal{P}$.
 Observe $y_{n+1} = f(\mathbf{x}_{n+1})$
 Increment n
end while
return datapoint $(\mathbf{x}, y) \in \mathcal{D}_N$ with the smallest evaluated f .

4.2.1 Gaussian process regression

The aim now is to use Eq. (4.19) to infer a probability distribution for the function f on the considered parameter space $\mathcal{P}$. The prior distribution of functions, which we base on the current set of data $\mathcal{D}_n$, is taken to be a GP.

A stochastic process is a collection of random variables that are indexed by some mathematical set, e.g. time, space or as in this case, the continuous space $\mathcal{X}$. By definition, a GP is a stochastic process such that every finite collection of random variables has a multivariate Gaussian distribution [54]. It can be thought of as the generalization of the Gaussian distribution over a finite vector space to a function space of infinite dimension. While a Gaussian is determined by its mean and covariance matrix, GPs are completely characterized by a mean function $\mu(\mathbf{x})$ and covariance function $k(\mathbf{x}, \mathbf{x}')$. Given $\mathcal{D}_n$, we indicate that the function f is modelled using a GP as

$$f(\mathbf{x})|\mathcal{D}_n \sim \mathcal{GP}(\mu(\mathbf{x}), k(\mathbf{x}, \mathbf{x}')). \quad (4.20)$$

For the real process $f(\mathbf{x})$, the mean and covariance functions are defined as the expectation values

$$\mu(\mathbf{x}) := \mathbb{E}[f(\mathbf{x})], \quad (4.21)$$

$$k(\mathbf{x}, \mathbf{x}') := \mathbb{E}[(f(\mathbf{x}) - \mu(\mathbf{x}))(f(\mathbf{x}') - \mu(\mathbf{x}'))] \quad (4.22)$$

For simplicity, we assume the prior mean function $\mu_0(\mathbf{x}) = 0$. Note however that this is not necessarily a strong limitation, since the mean $\mu_n(\mathbf{x})$ of the posterior process is not confined to be zero. Furthermore, if necessary, we can ensure the resulting approximation of the objective function is positive definite by performing the GP regression on $\log f$ rather than on f directly. The same is true for any input parameter $x_i \rightarrow \log x_i$.

The covariance function, also called the *kernel*, typically has the property that if $\|\mathbf{x} - \mathbf{x}'\| < \|\mathbf{x} - \mathbf{x}''\| \implies k(\mathbf{x}, \mathbf{x}') > k(\mathbf{x}, \mathbf{x}'')$, for some norm $\|\cdot\|$. In other words,

points that are closer to another in $\mathcal{X}$ are more strongly correlated and vice versa, which is of course only true if f is a continuous function. Additionally, kernels are positive semi-definite functions, i.e. $k(\mathbf{x}, \mathbf{x}') \geq 0, \forall \mathbf{x}, \mathbf{x}' \in \mathcal{X}$. It is also very common that they are stationary, meaning that $k(\mathbf{x}, \mathbf{x}') = k(\|\mathbf{x} - \mathbf{x}'\|)$. Here, we define this norm as

$$\|\mathbf{x} - \mathbf{x}'\|^2 = \sum_{i=1}^d \frac{(x_i - x'_i)^2}{\theta_i^2} =: r^2(\mathbf{x}, \mathbf{x}'; \boldsymbol{\Theta}), \quad (4.23)$$

where $\boldsymbol{\Theta} := \{\theta_i\}_{i=1}^d$ is a set of hyperparameters describing the correlation length in each input dimension. Each time when new data is sampled, any such hyperparameter within a GP is chosen such that the marginal likelihood $p(\mathcal{D}_n)$ in Eq. (4.19) is maximized. This procedure follows Algorithm 2.1 in *Gaussian Processes for Machine Learning* [54], using a Cholesky factorization for calculating the inverse of the covariance matrix, which appears in the update Eq. (4.28) and (4.29).

A commonly used covariance function is the *Gaussian kernel* and is given by

$$k_G(\mathbf{x}, \mathbf{x}'; \boldsymbol{\Theta}) := \exp\left(-\frac{1}{2}r^2(\mathbf{x}, \mathbf{x}'; \boldsymbol{\Theta})\right). \quad (4.24)$$

However, since this kernel is infinitely differentiable, the resulting GP is very smooth. In reality, this is often not the case for many physical processes and as the black-box function f is obtained from solving a strongly non-linear set of equations, this assumption may not be appropriate. Instead we make use of another covariance function, the *Matérn kernel*, given by

$$k_M(\mathbf{x}, \mathbf{x}'; \nu, \boldsymbol{\Theta}) := \frac{2^{1-\nu}}{\Gamma(\nu)} \left(\sqrt{2\nu} r^2(\mathbf{x}, \mathbf{x}'; \boldsymbol{\Theta})\right)^\nu K_\nu\left(\sqrt{2\nu} r^2(\mathbf{x}, \mathbf{x}'; \boldsymbol{\Theta})\right), \quad (4.25)$$

where Γ is the gamma function and K_ν is a modified Bessel function of the second kind. The additional hyperparameter ν sets the smoothness of the resulting function drawn from the GP and in the limit $\nu \rightarrow \infty$, the kernel $k_M(\mathbf{x}, \mathbf{x}') \rightarrow k_G(\mathbf{x}, \mathbf{x}')$. The Matérn covariance function in Eq. (4.25) is therefore considered as a generalization of the Gaussian kernel in Eq. (4.24). In this work we use the fixed value of $\nu = 5/2$.

To illustrate how the correlation length θ and the smoothness ν affect the distribution of functions, Fig. 4.2 shows five randomly sampled functions from three different GPs. The first two kernels, (a) and (b), both have a correlation length of $\theta = 1/2$ which yields a stronger correlation when compared to (c) with $\theta = 1/5$. This becomes apparent when comparing the heatmap of the corresponding covariance matrices, in which a strong correlation is indicated by bright colours. Matrix elements away from the highlighted diagonal more quickly decrease for case (c) compared to the other cases with larger correlation lengths, hence the larger variations in f . Furthermore, with a lower ν , the covariance matrix in case (a) is very peaked near the diagonal but does not decrease quite as rapidly as for case (c). This results in functions sampled from GPs using kernel (a) that are much less smooth compared to the other kernels, but it doesn't vary much at larger scales like in case (c).

To see how to incorporate the knowledge about $\mathcal{D}_n$ into making an informed prediction of the value of f at some new test point $\mathbf{x}_{n+1}$, a few matrices is constructed. Let $\mathbf{k}$ be a vector containing the covariance between $\mathbf{x}$ and each of the n points

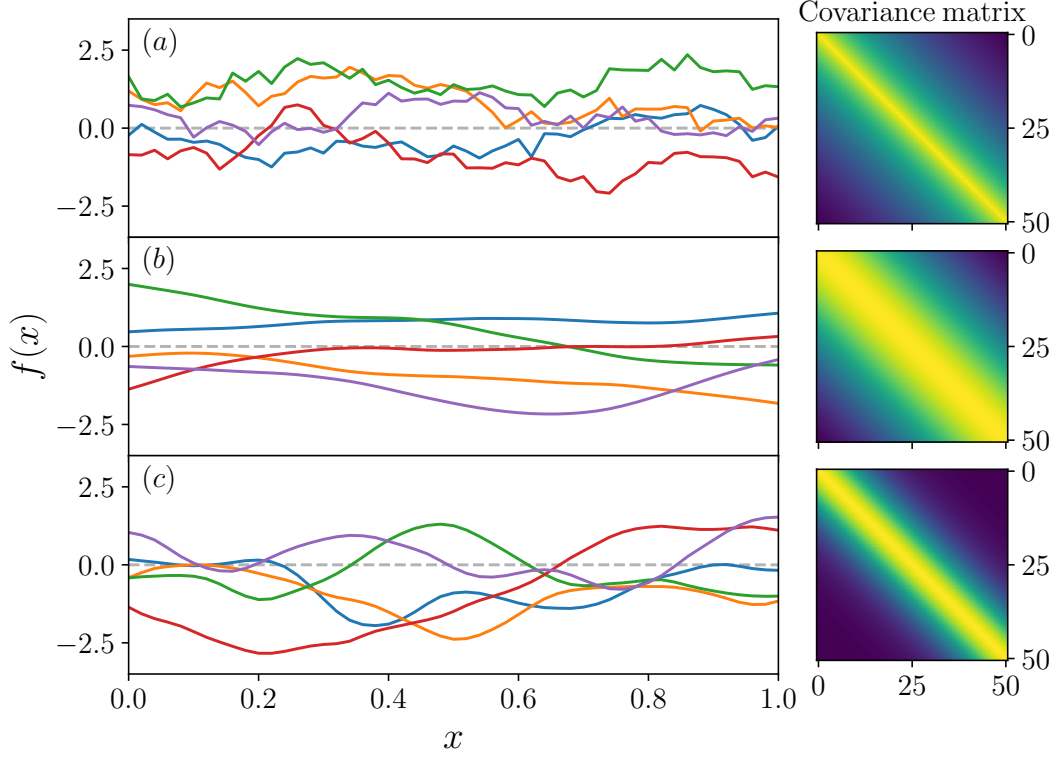


Figure 4.2: Sampled functions $f \sim \mathcal{GP}(0, k_M)$ using Matérn kernels in Eq. (4.25) with varying correlation lengths and smoothness θ and ν . Shown to the right is heatmaps of covariance matrices for the 50 uniformly placed inputs $x \in [0, 1]$. Each curve is made by connecting together the 50 corresponding inputs describing the multivariate Gaussian distributions, which used to generate five functions each. The kernel hyperparameters are (a) $\theta = 1/2$, $\nu = 1/2$, (b) $\theta = 1/2$, $\nu = 5/2$ and (c) $\theta = 1/5$, $\nu = 5/2$.

already sampled, $k_i := k(\mathbf{x}_{n+1}, \mathbf{x}_i)$, $i \leq n$, which we call the *test covariance vector*. Similarly, the *covariance matrix* Σ is constructed by letting its element correspond to the covariance between every sampled point in X_n , i.e. $\Sigma_{ij} := k(\mathbf{x}_i, \mathbf{x}_j)$, $i, j \leq n$. The set of objective function evaluations Y_n is turned into a vector $\mathbf{y}$, such that $y_i := f(\mathbf{x}_i)$. Consider adding the test point to the GP, which by definition is a collection of random variables with a distribution of a multivariate Gaussian. The joint distribution of the output data $\mathbf{y}$ and the test output y_{n+1} , given the set of sampled input data X_n and our prior assumption about the mean function $\mu_0(\mathbf{x})$ being zero, is thus given by

$$\begin{bmatrix} \mathbf{y} \\ y_{n+1} \end{bmatrix} \sim \mathcal{N} \left(\mathbf{0}, \begin{bmatrix} \Sigma + \sigma_\varepsilon^2 \mathbb{1} & \mathbf{k}^\top \\ \mathbf{k} & k(\mathbf{x}_{n+1}, \mathbf{x}_{n+1}) \end{bmatrix} \right), \quad (4.26)$$

where a Gaussian noise with variance σ_ε^2 is included, and $\mathbb{1}$ is the $n \times n$ identity matrix. To get the posterior distribution of functions, this joint distribution must be restricted to contain only those functions that agree with the previous evaluations,

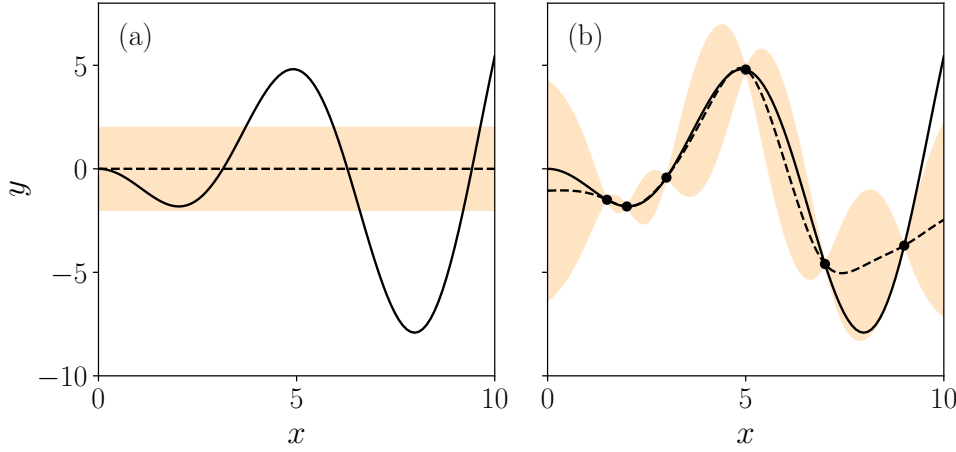


Figure 4.3: Both panels include the example objective function $f(x) = -x \sin x$ (solid black) on the interval $x \in [0, 10]$. The GP prior distribution is shown in panel (a), with a zero mean function (dashed black) including the 95% credible interval (yellow). Panel (b) includes six randomly sampled points (black dots) together with the corresponding posterior distribution. Both the mean and variance functions are updated by the Eqs. (4.28) and (4.29), respectively. Note that in regions further away from the sampled points, the GP is more uncertain about its prediction of f , hence the larger areas covered in yellow.

stored in $\mathcal{D}_n$. In probabilistic terms, this corresponds to *conditioning* the joint Gaussian prior distribution on $\mathcal{D}_n$ using Bayes' theorem from Eq. (4.19), which yields the conditional probability

$$y_{n+1} | \mathcal{D}_n, \mathbf{x}_{n+1} \sim \mathcal{N}(\mu_n(\mathbf{x}_{n+1}), \sigma_n^2(\mathbf{x}_{n+1})), \quad (4.27)$$

for which posterior mean and variance functions are calculated via

$$\mu_n(\mathbf{x}_{n+1}) := \mathbf{k}^\top [\Sigma + \sigma_\epsilon^2 \mathbf{1}]^{-1} \mathbf{y}, \quad (4.28)$$

$$\sigma_n^2(\mathbf{x}_{n+1}) := k(\mathbf{x}_{n+1}, \mathbf{x}_{n+1}) - \mathbf{k}^\top [\Sigma + \sigma_\epsilon^2 \mathbf{1}]^{-1} \mathbf{k}. \quad (4.29)$$

Algorithm 3 returns both the input $\mathbf{x}_n^*$ corresponding to the point with the smallest evaluated $y_i \in Y_n$, and also the minimizer of the posterior mean function $\mu_n(\mathbf{x})$, which given a sufficient number of samples in the vicinity of the actual optimum should yield $\sigma_n^2 \approx 0$ in this region, making μ_n an ever improving point estimate of the objective function f . As already mentioned in the beginning of this chapter, only deterministic and thus noise-free simulations are considered in this work, hence we omit the Gaussian noise by setting $\sigma_\epsilon^2 = 0$.

This procedure of updating the probability density for the estimate of the objective function f given some set of samples $\mathcal{D}_n$, is exemplified in Fig. 4.3 with a one-dimensional function. It shows how four function evaluations are incorporated into the posterior distribution of functions, shown to the right, given a prior $\mathcal{N}(0, 1)$, illustrated in the left panel. Note that the noise-free inference of this function yields

a mean function, shown in red, which is exactly equal to f in the sampled points, i.e. $\mu_n(\mathbf{x}_i) = f(\mathbf{x}_i)$.

4.2.2 Expected improvement acquisition function

Having established how knowledge about the function f is inferred from a set of sampled data $\mathcal{D}_n$ using GP regression in Sec. 4.2.1, this section describes a strategy for choosing an optimal new point $\mathbf{x}_{n+1}$ to sample and evaluate $y_{n+1} := f(\mathbf{x}_{n+1})$. The *expected improvement* $\text{EI}_n(\mathbf{x})$ is one of the most commonly used acquisition functions within Bayesian optimization, and can be derived from a thought experiment [52]. Let $f_n^* \in Y_n$ be the minimal value of f evaluated over all $\mathbf{x} \in X_n \subset \mathcal{P}$ that has currently been sampled, and let $\mathbf{x}_n^*$ be the corresponding input. If the optimization procedure is terminated at this point, Algorithm 3 would simply return $\mathbf{x}_n^*$ as the best estimate of the actual optimum $\mathbf{x}^*$.

Suppose that an additional evaluation is to be performed at any point $\mathbf{x}$ yielding $f(\mathbf{x})$. After this, the minimal observed value of f is either $f(\mathbf{x})$ if $f(\mathbf{x}) < f_n^*$ or still f_n^* otherwise. With the goal of finding the minimum of f , the improvement of such an evaluation in the former case is given by $f_n^* - f(\mathbf{x})$ and 0 in the latter. The best choice would be to maximize this improvement, but the actual value of $f(\mathbf{x})$ is still unknown before it being evaluated. Instead, what is done in Algorithm 3 is that the next sample is taken to be the maximizer of the expectation value of the improvement, given the trained GP. We can write this as

$$\mathbf{x}_{n+1} = \operatorname{argmax} \left(\mathbb{E}_n \left[\max(0, f_n^* - f(\mathbf{x})) \right] \right) =: \operatorname{argmax} \left(\text{EI}_n(\mathbf{x}) \right), \quad (4.30)$$

where $\mathbb{E}_n[\cdot]$ should be understood as the expectation under the posterior distribution, given the previously evaluated $\mathcal{D}_n$. Using the notation $\Delta_n(\mathbf{x}) := f_n^* - \mu_n(\mathbf{x})$ for the expected difference in quality between $\mathbf{x}$ and the currently best point, the expected improvement can be evaluated in closed form as

$$\text{EI}_n(\mathbf{x}) = \max(0, \Delta_n(\mathbf{x})) + \sigma_n \varphi \left(\frac{\Delta_n(\mathbf{x})}{\sigma_n(\mathbf{x})} \right) - |\Delta_n(\mathbf{x})| \Phi \left(\frac{\Delta_n(\mathbf{x})}{\sigma_n(\mathbf{x})} \right). \quad (4.31)$$

Above, $\varphi(\cdot)$ is the probability density of the standard normal distribution $\mathcal{N}(0, 1)$ and $\Phi(\cdot)$ is its cumulative probability function. These are given by

$$\varphi(t) := \frac{e^{-t^2/2}}{\sqrt{2\pi}}, \quad \Phi(t) := \int_{-\infty}^t ds \, \varphi(s). \quad (4.32)$$

With this strategy and selecting each new sample point $\mathbf{x}_{n+1}$ as in Eq. (4.30), the Bayesian optimizer sequentially gains more confidence about the global optimum $\mathbf{x}_n^* \rightarrow \mathbf{x}^*$, as more function evaluations are performed, much faster than by randomly sampling from a uniform distribution [55].

5

Optimization of disruption mitigation scenarios

Having described the two methods of black-box optimization we use, in Ch. 4, this chapter presents their application to various disruption mitigation scenarios. The general goal is to identify what densities of injected neon and deuterium produce the most favorable outcomes in a disruption mitigation, corresponding to the minimizer of an objective function f . Here, the main objective is to suppress the maximal runaway current $I_{\text{RE}}^{\text{max}}$ as much as possible under the constraint that the CQ time, τ_{CQ} , is bound between the two values in Eq. (2.42). These and the other relevant output variables $I_{\Omega}^{\text{final}}$ and η_{cond} , which we use to construct the objective functions, are defined in Sec. 2.5.2. We are not aiming to model the details of the material injection, instead we assume the material to be instantaneously deposited in the form of neutrals, uniformly distributed over the magnetic flux surfaces.

The evaluation of f consists of two steps: (1) running the disruption simulation and (2) calculating the value of f provided the simulation outputs. Step (1) involves using a standard multidimensional Newton's method for solving a set of nonlinear coupled equations, using the DREAM framework as described in Sec. 3.4. The approach of choice is to treat this first part as a black-box function, which we denote $b : \mathcal{X} \rightarrow \mathcal{Z}$. Here $\mathcal{X}$ and $\mathcal{Z}$ are the spaces of input parameters and output variables, respectively. The evaluation of $\mathbf{z} = b(\mathbf{x})$ is what by far takes the most time during the optimization procedures employed within this work. Although having the underlying system of equations described in Ch. 3, the black-box approach disregards any possible structures contained within this system of equations that could be exploited by using other more informed approaches of optimization. On the other hand, this allows for more general applications of the optimization framework to cases other than the considered fluid disruption simulations (e.g. ones also resolving the material injection process, or ones using kinetic electron modelling).

The second step (2) is to compute the value of the objective function from the simulation outputs. Letting the analytical function $\mathcal{L} : \mathcal{Z} \rightarrow \mathcal{Y} \subseteq \mathbb{R}_+$ be the objective function, the action of evaluating f at the point $\mathbf{x}$ is expressed as $f(\mathbf{x}) = \mathcal{L}(b(\mathbf{x}))$. The problem of optimizing a disruption simulation with respect to a set of input parameters reduces to localizing the global minimum of f .

In Sec. 5.1, the base objective $\mathcal{L}_1$ is defined in terms of the maximum RE current, the CQ time and the final Ohmic current. The disruption simulation considered in this section is initiated by prescribing an exponentially decaying temperature, which is subsequently evolved by the energy balance Eq. (3.9) after reaching temperatures below 100 eV, as described in Sec. 3.2.

In Sec. 5.2, the disruption is instead initiated by prescribing the heat transport, assuming different magnetic perturbation amplitudes $\delta B/B$. This allows us to self-consistently evolve the temperature, even during the TQ, which in turn makes it possible to include the transported heat fraction in the objective function $\mathcal{L}_2$. When the average temperature has dropped down below 100 eV, the transport is disabled. We perform optimizations similar to those in Sec. 5.1, for a range of values of the radially uniform magnetic perturbation $\delta B/B$. Selecting a value that yields an optimum which does not simultaneously satisfy the requirements put forth in Eq. (2.42), a four-dimensional optimization is then performed in Sec. 5.3, where the radial uniformity of the injected material is relaxed, in terms of two additional parameters that enter the optimization.

Finally in Sec. 5.4, we present a method of addressing the sensitivity of a given optimum by maximizing the objective function on smaller parameter regions centered around it. In doing so, this method provides an estimate of the worst possible outcome for input parameters in the vicinity of the optimum, representing an uncertainty in the amount of material deposited in the plasma.

5.1 Initially prescribed temperature evolution

When temperature evolution is prescribed during the thermal quench, the transported fraction of heat loss is not calculated, and as such, cannot be part of the objective function. Instead, the goal of the optimizations performed in this section is to minimize the sum of the maximum RE current and the final Ohmic current, under a constraint on the CQ time. The latter, $\theta(\tau_{\text{CQ}})$, is an indicator whether or not the disruption is to be rejected due to falling outside the tolerated range of CQ times, based on the decay rate of the plasma current. Namely, whenever $50 \text{ ms} < \tau_{\text{CQ}} < 150 \text{ ms}$, this metric is zero and otherwise it is set to a value of 100. However, since the GP regression of the objective function assumes f to be continuous in the Bayesian optimization, we refrain from defining $\theta(\tau_{\text{CQ}})$ in terms of the discontinuous Heaviside step function. Instead, we make use of the differentiable logistic function $\ell(x) = [1 + \tanh(x/2)]/2$ and let

$$\theta(\tau_{\text{CQ}}) := 100\ell(k(50 - \tau_{\text{CQ}} [\text{ms}])) + 100\ell(k(\tau_{\text{CQ}} [\text{ms}] - 150)), \quad (5.1)$$

where $k = 300$ controls how rapidly the function rises at the interval bounds. Replacing the logistic function with the Heaviside step function thus corresponds to Eq. (5.1) in the limit $k \rightarrow \infty$. Note that if the black-box function b contains a discontinuity, such a feature will propagate into the objective function. Nevertheless, this definition of $\theta(\tau_{\text{CQ}})$ will not impose any further discontinuities.

This leaves the two remaining output variables, $I_{\text{RE}}^{\text{max}}$ and $I_{\Omega}^{\text{final}}$, to be examined. Figure 5.1 displays the resulting landscapes of these two currents in the parameter space of injected material quantities, based on the outputs of 1600 simulations. These sampled points are uniformly distributed on logarithmic scales on $\mathcal{P}_1$, which is spanned by the densities of injected deuterium $n_{\text{D}} \in [10^{18}, 1.6 \cdot 10^{22}] \text{ m}^{-3}$ and neon $n_{\text{Ne}} \in [10^{15}, 10^{20}] \text{ m}^{-3}$. The reason for this particular value for the upper bound in n_{D} is a practical one; that the solver in DREAM has trouble converging for

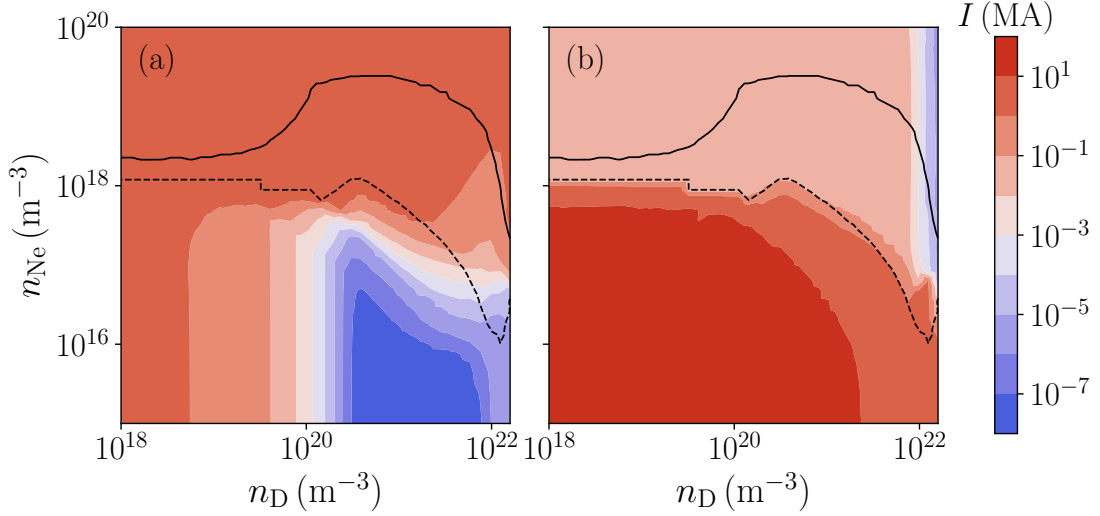


Figure 5.1: The maximum RE current (a) and final Ohmic current (b) as functions of the injected deuterium and neon densities, n_D and n_{Ne} . Above the dashed (solid) line, the CQ time is less than 150 ms (50 ms).

$n_D > 1.6 \cdot 10^{22} \text{ m}^{-3}$, due to the resulting high density and low temperature leading to very rapid recombination, even if the time resolution is considerably increased.

Recalling that the avalanche growth rate in Eq. (3.19) increases with the ration of the total electron density (free+bound) and the free electron density, this explains the high runaway currents observed in Fig. 5.1a for large amounts of injected neon.

We construct the base objective function $\mathcal{L}_1$ as a weighted sum of the three metrics used here, as follows

$$\mathcal{L}_1(I_{\text{RE}}^{\text{max}}, I_{\Omega}^{\text{final}}, \tau_{\text{CQ}}) := \frac{I_{\text{RE}}^{\text{max}}}{I_{\text{RE}}^{\text{crit}}} + \frac{I_{\Omega}^{\text{final}}}{I_{\Omega}^{\text{crit}}} + \theta(\tau_{\text{CQ}}). \quad (5.2)$$

Here, we use the parameters $I_{\text{RE}}^{\text{crit}} = 150 \text{ kA}$ and $I_{\Omega}^{\text{crit}} = 300 \text{ kA}$, for which the former is set in accordance with the mitigation requirements set forth by the ITER mitigation task force [38].

If the metric for the final Ohmic current were to be excluded from the objective function, Fig. 5.1a indicates that the optimum then would be located roughly at $n_D \sim 10^{22}$ and $n_{\text{Ne}} \sim 10^{16} \text{ m}^{-3}$ (i.e. around the lowermost point of the dashed curve), where $I_{\text{RE}}^{\text{max}}$ is at its minimum under the constraint concerning τ_{CQ} . Comparing this to Fig. 5.1b, this region in parameter space corresponds to disruptions with substantial values of $I_{\Omega}^{\text{final}}$. Further analysis of individual simulations in this region reveals the presence of one or several radial spikes of concentrated Ohmic current density $j_{\Omega}(r)$ that materialize inside the plasma. Similar soliton-like spikes have previously been observed in various numerical disruption scenarios [56, 57]. In practice, turbulence in the plasma is believed to eventually cause any such spikes to dissolve before gaining too much current.

Nevertheless, the desired mitigation outcome is the suppression of both currents, considering the possibility of significant portions of any lingering Ohmic currents

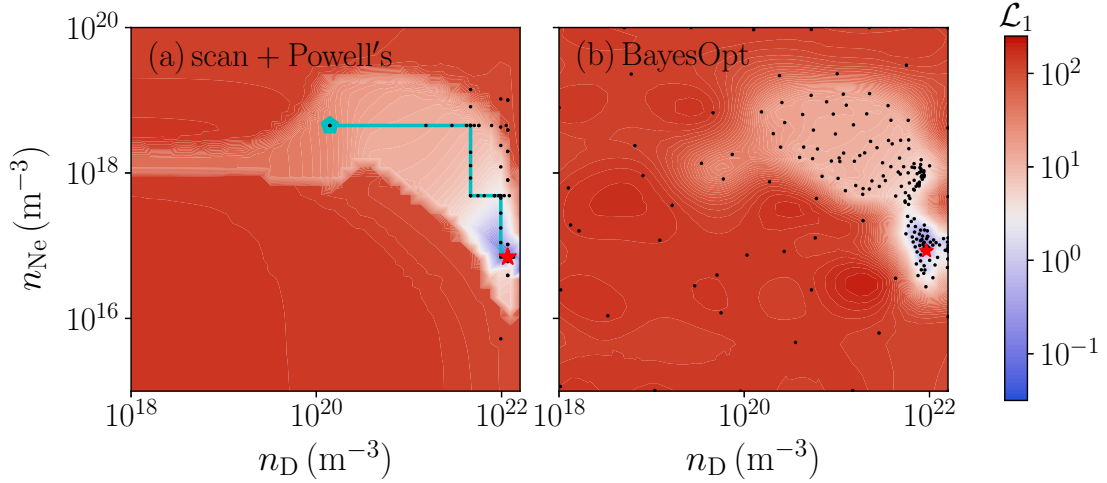


Figure 5.2: The objective $\mathcal{L}_1$ as functions of the input parameters of injected deuterium n_D and neon n_{Ne} . Panel (a) shows a scan of function evaluations on a logarithmically equidistant grid, of 1600 samples on $\mathcal{P}_1$, as well as the trajectory (cyan lines) and corresponding 118 function evaluations (black points) from an optimization using Powell’s method. Panel (b) shows is the posterior mean function of a GP, which has been trained on 200 evaluations. The found optima are indicated in both figures (red stars).

transforming into runaways, as motivated in Sec. 2.5.2. A conservative study should account for this, as the fate of such Ohmic channels cannot be inferred from our simulation. By penalizing such scenarios in which $I_{\Omega}^{\text{final}}$ is non-negligible, we intend to avoid such scenarios. Assuming a 50% conversion, we arrive at $I_{\Omega}^{\text{crit}} = 2I_{\text{RE}}^{\text{crit}} = 300 \text{ kA}$.

Using the objective function $\mathcal{L}_1$ defined in Eq. (5.2), we perform two individual optimizations on $\mathcal{P}_1$, using Powell’s and the Bayesian methods, respectively. Fig. 5.2a illustrates the resulting landscape of $\mathcal{L}_1$ based on the scan shown in Fig. 5.1 and the path taken by Powell’s method. The first point on its trajectory is indicated by the cyan dot with the coordinates $n_D = 1.4 \cdot 10^{20}$ and $n_{Ne} = 4.5 \cdot 10^{18} \text{ m}^{-3}$, which was chosen pseudo-randomly with the condition $50 \text{ ms} < \tau_{CQ} < 150 \text{ ms}$. After a total of 118 function evaluations, Powell’s method converges to a minimum at $n_D = 9.27 \cdot 10^{21}$ and $n_{Ne} = 8.51 \cdot 10^{16} \text{ m}^{-3}$, with $\mathcal{L}_1 = 0.021$ (indicated by the red star). The corresponding time evolution of the RE and Ohmic currents at the initial and optimal point are shown in Fig. 5.3. Comparing how the two sets of output variables change, the maximum RE current reduces from 4000 to 2.0 kA, the final Ohmic current from 26 to 2.2 kA and the τ_{CQ} increases from 69.8 to 82.8 ms.

Furthermore, the Bayesian optimization is displayed in Fig. 5.2b, which shows the posterior mean function of a GP based on 200 function evaluation. We see that the found optimum is not located at exactly the same parameters as found by Powell’s method. It finds $n_D = 1.17 \cdot 10^{22}$ and $n_{Ne} = 7.00 \cdot 10^{16} \text{ m}^{-3}$, with the corresponding $\mathcal{L}_1 = 0.078$, being roughly four times higher than Powell’s. However, since $\mathcal{L}_1 < 1$ in both cases, these differences are of no concern since both meet the mitigation criteria with good margins. Due to the multiple orders of magnitude of injected

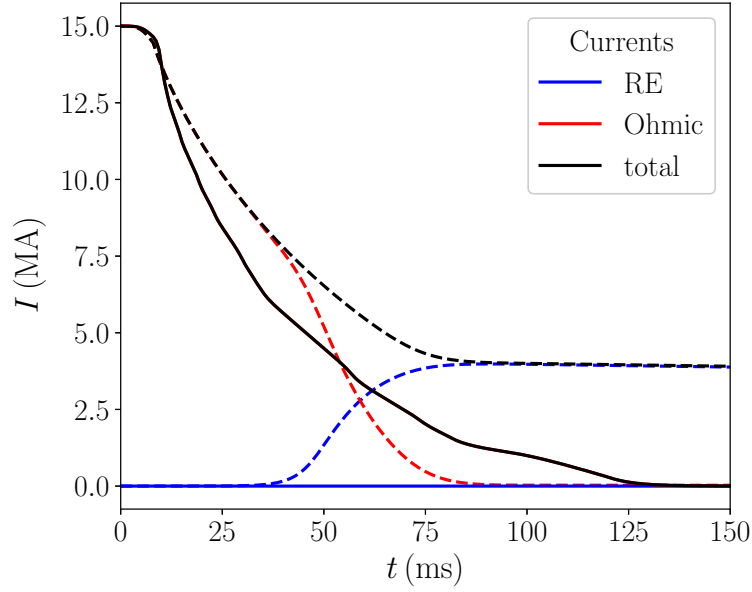


Figure 5.3: Time evolution of the Ohmic and RE currents corresponding to the first point sampled (dashed lines) in Powell’s method and that of the found optimum (solid lines). The corresponding two sets of parameters are indicated in Fig. 5.2a.

materials covered by $\mathcal{P}_1$, we optimize over the base-10 logarithm of the injected densities in the Bayesian method. After 200 evaluations, the obtained correlation lengths in $\log(n_D [\text{m}^{-3}])$ and $\log(n_{\text{Ne}} [\text{m}^{-3}])$ are 0.581 and 0.471, respectively.

After only 100 evaluations, the GP has not yet sampled this region where the optimum is found. At this stage, it uses the sparsely sampled points around the true optimum and interpolates the value of the objective function in this region, resulting in a prediction of $\mathcal{L}_1$ that is much higher than it actually is. The posterior mean function at various total number of evaluations is included in Appendix B.1, displaying the evolution of the Bayesian optimization.

On a final note, it should be emphasized that both of the obtained optima satisfy the constraints given in Eq. (2.42). It appears that a relatively high amount of injected deuterium, in combination with a relatively low amount of injected neon results in a heavily suppressed generation of REs. This is expected since the lighter deuterium isotopes are mostly fully ionized, corresponding to a large amount of free electrons that increase the drag force experienced by the REs, whereas the neon can be partially ionized, entailing less friction due to the screening effect mentioned in Sec. 2.5.2.1. Furthermore, it is interesting to note the adequate CQ time. With less neon to radiate away energy at lower temperatures, the resulting CQ time would in general be longer. However, since an increase in the injected deuterium results in a higher free electron density, implying that there are more electrons available to excite the neon atoms, the radiated power increases as a result. In addition, at sufficiently low temperatures the deuterium can recombine as well, contributing to the total radiated heat losses. As such, with a sufficiently large amount of injected deuterium, the temperature drop can be rapid enough to obtain a CQ time of under 150 ms. Another effect of the line radiation not taken into account in this

optimization is its ability to reduce the transported heat losses. This will be the next object of study.

5.2 Prescribed heat transport

As previously mentioned in Ch. 2, a major cause for the sudden loss in thermal energy during the TQ is radial transport due to magnetic perturbations induced by MHD instabilities. The physics behind this constitutes its own area of research and its self-consistent modelling is beyond the scope of the DREAM-code [58]. In order to study the conducted thermal losses we therefore adopt a simple model of the phenomena, where heat transport, as well as transport of REs, is prescribed by the Rechester-Rosenbluth diffusion coefficient from Eq. (3.10). It is assumed that the perturbation strength, $\delta B/B$, has a uniform spatial profile and remains constant throughout the TQ until the temperature has decreased to 20 eV, after which it is turned off. Using this method the temperature can be evolved self-consistently during the entire disruption simulation, taking into account additional factors such as Ohmic heating of the inductive system, collisions between electrons and different ion species, and line radiation from neutral and partially ionized isotopes. As discussed in Sec. 3.2, this entails a more complex TQ process that will be heavily dependent on the amount of injected material. Because of this, it is important that the temperature evolution is analyzed, making sure that the temperature has sufficiently decreased in order to represent a characteristic disruption scenario. If the final mean temperature does not fall below 0.1 % of the initial mean temperature, the simulation is penalized by letting the objective function return a large value.

Having a self-consistent temperature evolution now allows us to compute the amount of heat that is transported out from the plasma, constituting to localized heat loads on plasma facing components. The objective function is therefore updated to include a metric for the corresponding quantity, presented in Eq. (2.41):

$$\mathcal{L}_2(I_{\text{RE}}^{\text{max}}, I_{\Omega}^{\text{final}}, \tau_{\text{CQ}}, \eta_{\text{cond}}) := \mathcal{L}_1(I_{\text{RE}}^{\text{max}}, I_{\Omega}^{\text{final}}, \tau_{\text{CQ}}) + 10 \frac{\eta_{\text{cond}}}{\eta_{\text{cond}}^{\text{crit}}}. \quad (5.3)$$

Here the critical value of the conducted fraction is $\eta_{\text{cond}}^{\text{crit}} = 10\%$, and the term has been weighted by a factor 10 to increase its significance during optimization.

Figure 5.4 shows the posterior mean function of the GPs resulting from Bayesian optimization of simulations with heat transport with respect to the amount of injected deuterium and neon, using the objective function given in Eq. (5.3). The optimization was carried out for the different perturbation amplitudes $\delta B/B = 0.2\%, 0.3\%, 0.4\%, 0.5\%$, and a total of 110 sampled points were used. Furthermore, the penalty for an insufficient TQ was set to 500. Powell's method was also applied to the different scenarios, using an initial starting point $n_{\text{D}} = 6 \cdot 10^{20} \text{ m}^{-3}$, $n_{\text{Ne}} = 6 \cdot 10^{17} \text{ m}^{-3}$, and its path is illustrated in the figure as well.

One apparent feature of the contours shown in Fig. 5.4 is that they only decrease significantly within a small region of the parameter space (the majority of the surface is in the red part of the spectrum). This is mostly due to the penalization of improper TQs. As mentioned in Sec. 3.2, the transported power density decreases at lower temperatures, while the radiated power density generally increases. In order to

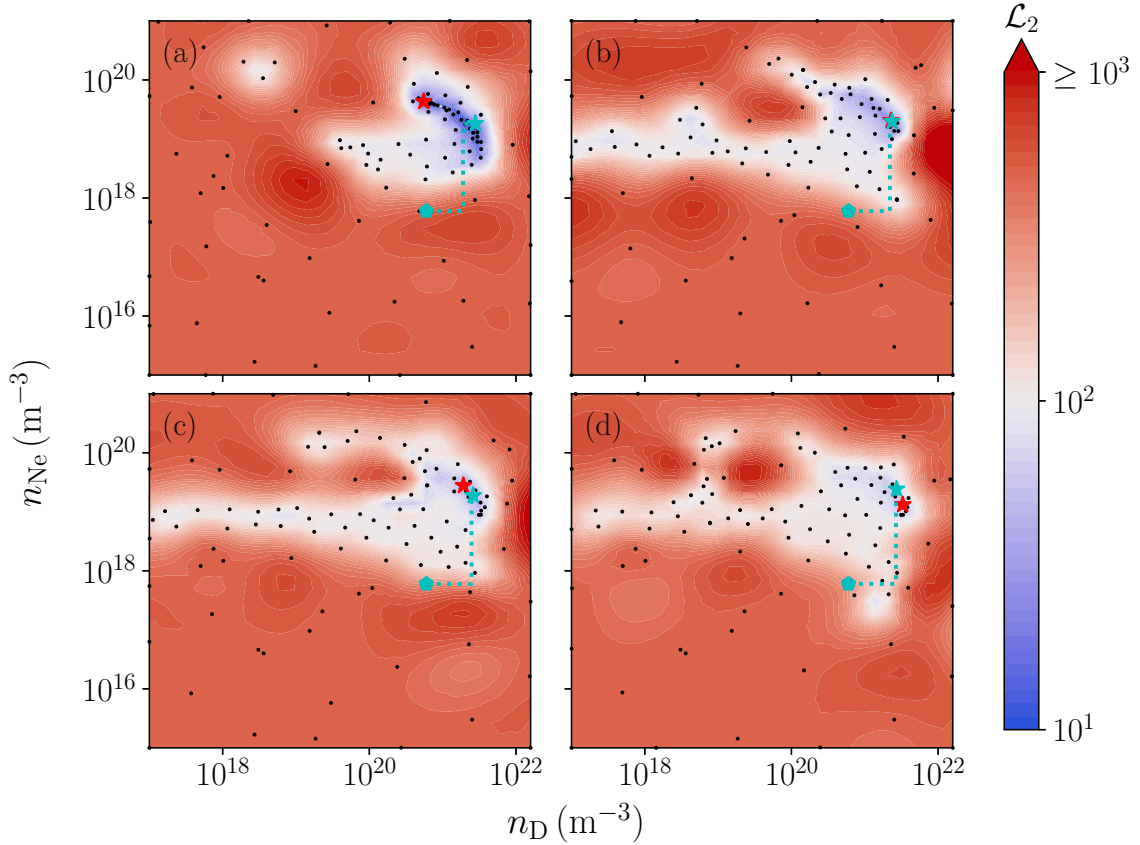


Figure 5.4: Heatmap of the posterior mean function of GPs obtained from Bayesian optimization of transport simulations in DREAM. The black dots indicate the points sampled during the process. The simulations assumed a uniform magnetic perturbation, $\delta B/B$, of (a): 0.2%, (b): 0.3%, (c): 0.4%, (d): 0.5%, throughout the TQ. The location of the minimal sampled point is denoted by a red star, while the starting point, path, and optimum of Powell’s method are indicated by a cyan pentagon, dotted line, and star respectively.

obtain a full TQ a radiative collapse is needed to exhaust the remaining heat. It is therefore not surprising that the bottom regions in all four cases show a constant, high objective function, as the amount of neon in these simulations is not sufficient to efficiently radiate away the residual thermal energy. The limit is located somewhere below $n_{\text{Ne}} = 10^{19} \text{ m}^{-3}$, with a slight decrease towards higher deuterium densities. This decrease could potentially be a consequence of the increased dilution entailing a lower plasma temperature, in combination with the increased free electron density, resulting in amplified radiative heat losses.

Additionally, one may note that the overall size of these admissible regions decreases for lower perturbation strengths. This is most evident in the $\delta B/B = 0.2\%$ scenario where the “valley”, usually located around a neon density of $n_{\text{Ne}} = 10^{19} \text{ m}^{-3}$, has disappeared and a complete TQ is not obtained for the previously adequate neon densities, unless a sufficiently high deuterium density, around 10^{19} m^{-3} or higher, is injected. After the assimilation of the initially injected material, the temperature evolution is generally dominated by heat transport until it is sufficiently low

to allow for a significant recombination of the neon. It is possible that reducing the transport even further could result in this transition never occurring, effectively prohibiting a radiative collapse. This idea is supported by the fact that the problem is resolved for higher deuterium densities, since injecting more material causes the initial temperature to drop faster through dilution and ionization.

As a final remark regarding the TQ, there is an unexpected non-monotonic behavior with respect to the neon density. In parts of the region where $10^{19} \text{ m}^{-3} < n_{\text{Ne}} < 10^{20} \text{ m}^{-3}$, there is yet again an insufficient temperature decay. This phenomenon is most apparent in the $\delta B/B = 0.4\%$ case in Fig. 5.4. Such behavior is surprising, since there would be further ionization and more material that could radiate away heat in this scenario. The studied system is however strongly non-linear and it is possible that the change in effective charge number, Z_{eff} , which affects the Ohmic heating, along with interactions on the atomic level, for instance ionization states and the power spectrum of the line radiation, may yield this result.

As for the non-penalized simulations with an adequate temperature decay, the majority gave values of the objective function on the order of $\mathcal{L}_2 \sim 10^2$. The most favorable outcomes for all the different perturbation levels is obtained in the vicinity of $n_{\text{D}} \sim 10^{21} \text{ m}^{-3}$ and $n_{\text{Ne}} \sim 10^{19} \text{ m}^{-3}$. The locations of the computed minima, summarized in Tab. 5.2, are relatively close to each other for most cases, and the two optimization techniques agree fairly well. The exception to this is the result for $\delta B/B = 0.2\%$, where the optimal parameters found using Bayesian optimization have shifted notably towards a lower deuterium density and higher neon density. Powell's method does however return an optimum that is closer to that of the other scenarios. The objective function and output values at the optima obtained from Bayesian optimization, as well as Powell's method, are presented in Tab. 5.1. Comparing the objective function at the location returned by the two methods for $\delta B/B = 0.2\%$, it is seen that Powell's method converged towards a local minimum where $\mathcal{L}_2 = 7.77$, while the Bayesian optimization was able to find a slightly lower minimum for which $\mathcal{L}_2 = 6.32$, that may represent the global minimum.

The optimum value of the objective function obtained by both methods increases significantly at higher magnetic perturbation strengths. To understand why, one must look at the individual outputs found in Tab. 5.1. One interesting feature of the results is that all output values, aside from the transported fraction, are within the acceptable limit in all the different scenarios. The maximum runaway current is below the critical value $I_{\text{RE}}^{\text{crit}} = 150 \text{ kA}$, by a good margin at each optimum.

Considering the parameters at which this optimum is found and comparing them to the scan in Fig. 5.1, where a prescribed exponential temperature decay was used, one might expect a much larger runaway current, especially considering the high neon densities, over 10^{19} m^{-3} . It is however important to note the significance of the runaway diffusion, introduced in Sec. 3.3, present in the simulations with transport. A sufficiently high transport of the REs during the TQ can severely limit the seed current. As such, it is possible that the total runaway current does not reach critical magnitudes, even when amplified by the avalanche mechanism. This offers an explanation as to why the optima appear in a previously unsuitable region of the parameter space. Despite the mentioned suppression, it should be noted

Table 5.1: Performance indicators of the optima in simulations with transport, found using both Bayesian optimization and Powell’s method.

Method	$\frac{\delta B}{B}$ (%)	$\mathcal{L}$	$I_{\text{RE}}^{\text{max}}$ (kA)	$I_{\Omega}^{\text{final}}$ (kA)	τ_{CQ} (ms)	η_{cond} (%)
Bayesian	0.2	6.32	$6.55 \cdot 10^{-5}$	$2.71 \cdot 10^{-9}$	57.3	5.1
	0.3	19.2	$1.41 \cdot 10^{-5}$	$2.58 \cdot 10^{-13}$	55.5	16
	0.4	37.5	$2.88 \cdot 10^{-2}$	$8.94 \cdot 10^{-14}$	53.4	26
	0.5	55.6	$6.26 \cdot 10^{-3}$	$2.66 \cdot 10^{-15}$	54.6	50
Powell	0.2	7.77	$6.61 \cdot 10^{-7}$	$1.29 \cdot 10^{-13}$	55.3	3.8
	0.3	19.2	$9.83 \cdot 10^{-6}$	$2.06 \cdot 10^{-13}$	55.5	16
	0.4	35.7	$1.59 \cdot 10^{-3}$	$2.48 \cdot 10^{-14}$	54.5	30
	0.5	55.2	$1.23 \cdot 10^{-1}$	$9.97 \cdot 10^{-16}$	52.3	35

that the maximum runaway current increases by several orders of magnitude for the higher perturbations in Table 5.1. It is likely that this is an effect of some runaway seed becoming more prominent during the TQ. For example, at the optima obtained using Powell’s method for $\delta B/B = 0.2\%$, the TQ occurs on a timescale of $\sim 1.8\text{ms}$, while an equivalent temperature drop only takes $\sim 0.5\text{ms}$ at the optima for $\delta B/B = 0.5\%$. As discussed in Sec. 3.3, the Dreicer and hot-tail runaway mechanisms are both sensitive to the temperature evolution inside the plasma and it is possible that an increased heat transport results in a subsequent amplification of these effects. Furthermore, it is possible that the RE transport redistributes the runaway seed in such a way that the avalanche generation becomes more efficient.

Moving on to the final Ohmic current, it is found to be vanishingly small for all optima listed in Tab. 5.1. This is expected since, as discussed in Sec. 5.1, the formation of Ohmic channels (or spikes) is counteracted by a sufficiently high heat diffusivity. As such, it seems that this phenomenon does not pose the same issue as it did in the exponential temperature decay scenario: It is apparent that neither the runaway nor Ohmic currents have a substantial impact on the objective function at any of the optima in question. The main contributions are instead the CQ time and conducted heat losses.

While τ_{CQ} lies within the acceptable interval $50 - 150\text{ms}$, the values at all optima are relatively close to the lower limit, implying that the corresponding contribution to the objective function at these points is no longer negligible (recall that the transition of the τ_{CQ} penalty across the threshold values is narrow, but smooth, giving finite values close to the threshold). Additionally the transported fraction, while tolerable for the $\delta B/B = 0.2\%$ case, undergoes a steep increase as the perturbation strength is increased. The transported power density grows monotonically with respect to the diffusion coefficient, which has a quadratic dependence on the magnetic perturbation amplitude. It is therefore expected that the radiated power density will eventually not be sufficient to avoid substantial transported heat loss, unless a very high amount of neon is injected. As previously pointed out, simply increasing the neon density can entail both an amplified runaway generation as well as a reduced CQ time, the latter being the main concern at these optima. It appears that these competing factors are what place the optima in this particular region.

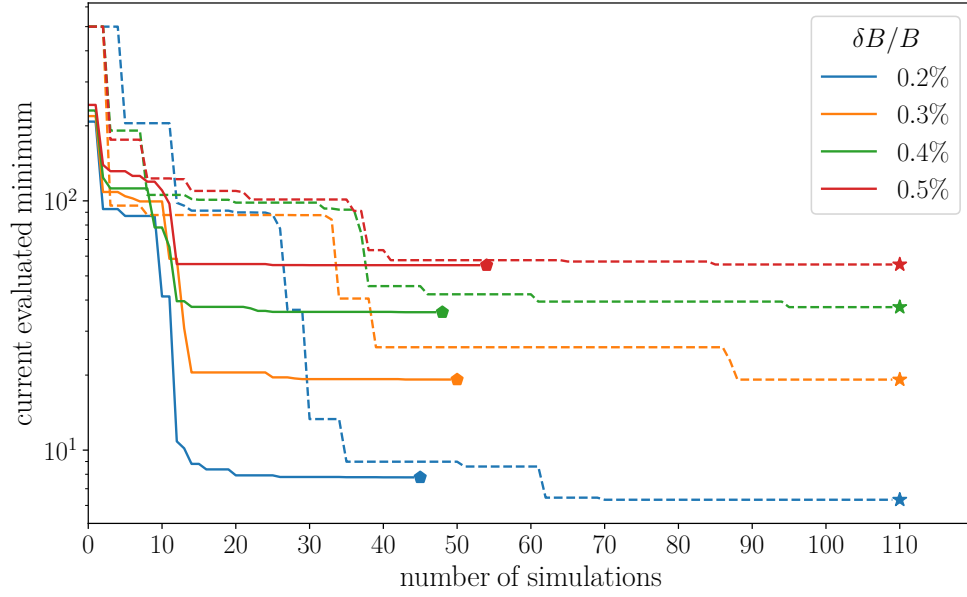


Figure 5.5: The lowest computed value during optimization with Powell’s method (solid) and the Bayesian method (dashed) as a function of the number of simulations performed. While the Bayesian optimization was carried out for a fixed number of 110 samples, Powell’s method is aborted once a convergence condition is satisfied.

Concluding this discussion, it should be noted that predicaments were found for both low and high magnetic perturbations, since the former could entail difficulties with incomplete TQs while the latter yielded a substantial transported fraction.

Besides studying the physical properties of disruption scenarios in different parameter regions, it is interesting to also analyze the results from an optimization standpoint, comparing the performance of the two methods. It was stated earlier that both methods converged towards similar minima, as could be seen in Fig. 5.4. However, looking at the figure one could argue that the starting point used in Powell’s method is rather advantageous. It should be noted that different starting points were employed during the testing stage of the method and, while many yielded an optimum in the same region, in some cases the routine returned parameter values differing by several orders of magnitude, indicative of a local minimum being found. Furthermore, if the routine is initialized inside of the penalized region, then the constant value returned by the objective function can make it impossible for Powell’s method to progress. Bayesian optimization does not run into the same problems since it scans the domain globally.

Another important aspect is the computational cost of using the respective methods. Figure 5.5 shows the current lowest evaluation of the objective function at a given number of simulations during optimization of the transport scenarios using both Powell’s method and Bayesian optimization. It is clear that Powell’s method converges significantly faster towards a minimum. This is not so surprising, since the acquisition function used when sampling with the Bayesian method favors “exploration” of the full parameter domain. Recall that in the case where $\delta B/B = 0.2\%$, the Bayesian optimization routine is able to find a better minimum than Powell’s,

as the latter gets stuck in a local optimum.

Lastly, a crucial difference between the two techniques is the informational gain. While Powell's method was under the right conditions able to efficiently localize an optimum, the Bayesian construction of a posterior probability distribution yielded the function maps seen in Fig. 5.5. These provide an advantageous insight into the behavior of the function that not only can be used for optimization purposes, but also makes it possible to discover new relations, such as the non-monotonic behavior of the TQ with respect to the injected neon density that was discussed earlier.

5.3 Profile optimization

In the previous section it was found that optimization based on the injected material densities alone may not yield an adequate disruption scenario if the magnetic perturbation strength is high during the TQ. It was hypothesized that the large diffusive transport needed to be compensated for by an increased heat radiation which, in turn, could lead to more efficient runaway generation and lower CQ times. These results motivate an extension of the optimization space to a higher dimensionality. The new parameters would suitably affect the temperature evolution inside the plasma, altering the balance between the different power densities. For this purpose, an additional degree of freedom was added to the injected material densities, allowing for variation of the spatial profile. A model is assumed where the density, $n(r)$, can be expressed as a function of the minor radius based on a profile, $g(r)$, according to

$$g(r; c) = 1 + \tanh \left[c \left(\frac{r}{a} - \frac{1}{2} \right) \right], \quad (5.4)$$

$$n(r; \bar{n}, c) = \bar{n} g(r; c) \frac{\int_0^a V' dr'}{\int_0^a V' dr' g(r'; c)}. \quad (5.5)$$

Here a denotes the plasma minor radius, $\bar{n}$ is the average density, V' is a spatial Jacobian for the r -coordinate, defined in Eq. (A.4) of Appendix A, and $c \in \mathbb{R}$ is a parameter governing the radial variation of the density. Note that $c < 0$ corresponds to a higher concentration toward the plasma center, while $c > 0$ implies that more material is deposited near the plasma edge. $c = 0$ is equivalent to a uniform distribution, similar to before, with density $n(r) = \bar{n}$. Equation (5.5) assures that the total number of particles, given by $N = \bar{n}V$, is held fixed for different profiles, i.e. different values of c .

With this model the injected material density profiles, $n_D = n_D(r; \bar{n}_D, c_D)$ and $n_{Ne} = n_{Ne}(r; \bar{n}_{Ne}, c_{Ne})$, are now described by four parameters which are the targets of optimization. Adopting this new parameter space, Powell's method and Bayesian optimization is subsequently applied to the worst performing transport scenario, for which $\delta B/B = 0.5\%$. Powell's method was given the initial starting point $\bar{n}_D = 6 \cdot 10^{20} \text{ m}^{-3}$, $c_D = 0$, $\bar{n}_{Ne} = 6 \cdot 10^{17} \text{ m}^{-3}$, $c_{Ne} = 0$, and the Bayesian routine used a total of 525 samples. Parameter values were in both cases confined to the intervals: $\bar{n}_D \in [10^{17}, 10^{22}] \text{ m}^{-3}$, $c_D \in [-12, 12]$, $\bar{n}_{ne} \in [10^{15}, 10^{21}] \text{ m}^{-3}$, $c_{Ne} \in [-12, 12]$. The

resulting optima for two methods is included in the summary found in Tab. 5.2. Note that for simplicity, the bar is removed when denoting the average densities in the table. The minimum returned by Powell's method yielded $\mathcal{L} = 9.65$, for which $I_{\text{RE}}^{\text{max}} = 55.4 \text{ kA}$, $\tau_{\text{CQ}} = 67.33 \text{ ms}$, and $\eta_{\text{cond}} = 9.28 \%$. For the Bayesian optimization the optima was even better, with a resulting objective function $\mathcal{L} = 5.50$, where $I_{\text{RE}}^{\text{max}} = 4.26 \text{ kA}$, $\tau_{\text{CQ}} = 63.87 \text{ ms}$, and $\eta_{\text{cond}} = 5.45 \%$. In both cases the remaining Ohmic current was once again very small, on the order of $\sim 10^{-11} \text{ kA}$ and $\sim 10^{-8} \text{ kA}$ for Powell's and the Bayesian method respectively.

It is important to point out that both optima performed significantly better than that of the analogous scenario with uniform injected material densities in Tab. 5.1 (reaching $\mathcal{L} \approx 55$). All of the presented output values lie within their respective acceptable ranges, with the transported fraction being significantly smaller and the current quench time having a substantial margin with respect to the lower limit. The location of the optima found by the two methods are however quite different and the distinctly smaller objective function returned by the Bayesian optimization suggests that it might correspond to a global optimum, whereas Powell's method converged toward a local optimum. Such behavior is expected to become more frequent as the dimensionality increases and the function becomes more complex.

The optimal injected density profiles are shown in panel a of Fig. 5.6. The deuterium and neon densities have been normalized with respect to the optimal uniform densities given by Powell's method in Sec. 5.2 ($\tilde{n}_\alpha = n_\alpha(r)/n_\alpha^{\text{uniform}}$). A common feature of the two results is that $\bar{n}_{\text{Ne}}^*$ has increased compared to the previous optima, with the vast majority being deposited near the plasma edge. The radiated power density is therefore higher further out in the plasma, which appears to be a valid strategy for avoiding considerable transported heat loss while avoiding extreme temperature drops in the center where the current density is higher and avalanche generation is more prominent. Furthermore, a more benign temperature evolution in the region where the Ohmic current density is the largest may also entail a longer current quench time.

At the minimum obtained using Bayesian optimization, the average neon density, $\bar{n}_{\text{Ne}}^* = 1.08 \cdot 10^{20} \text{ m}^{-3}$, is almost twice as high as that obtained using Powell's method, $\bar{n}_{\text{Ne}}^* = 5.24 \cdot 10^{19} \text{ m}^{-3}$. Furthermore, the shaping parameters, $c_{\text{Ne}}^* = 6.3$ and $c_{\text{Ne}}^* = 10$ respectively, imply a slightly smoother profile for the Bayesian optimum. As such, the latter has a notably higher neon density closer to the center of the plasma than that found by Powell's method. While this might counteract transported heat losses, based on the previous reasoning this could also lead to a faster CQ time and amplified runaway generation. Indeed the CQ time is lower than the one obtained at the other optimum, but only slightly, and the maximal runaway current is in fact smaller by an order of magnitude. A possible explanation as to why a larger, more evenly distributed neon injection did not lead to further complications regarding current evolution and runaway generation can be found in the significantly different deuterium densities returned by the two methods. Looking at Tab. 5.2, the Powell optimum has an average density, $\bar{n}_{\text{D}}^* = 4.1 \cdot 10^{21} \text{ m}^{-3}$, that is higher than that of the uniform scenario along with a positive, although subtle, inclination because $c_{\text{D}}^* = 0.44$. In contrast, the Bayesian optima has an average density, $\bar{n}_{\text{D}}^* = 3.57 \cdot 10^{20} \text{ m}^{-3}$, that is notably lower and a profile that is much more peaked near the center, since

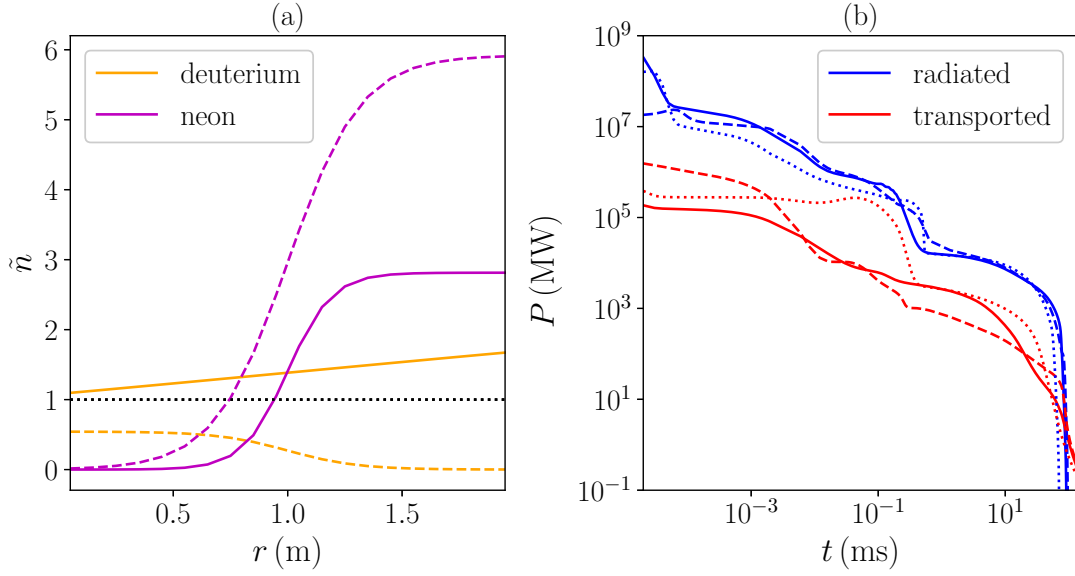


Figure 5.6: Panel (a): Injected density profiles from the optimal disruption scenario including heat transport with $\delta B/B = 0.5\%$, obtained using Powell's method (solid) and Bayesian optimization (dashed) on a four-dimensional parameter space. The deuterium and neon profiles are normalized to a baseline case corresponding to the values found by the Powell's optimization using constant densities. Panel (b): Evolution of the radiated and transported power during simulations using the corresponding density profiles of the optima, as well as the uniform profiles of the baseline case (dotted).

$c_D^* = -6.6$. The lower amount of injected deuterium reduces the initial temperature decrease caused by ionization and dilution and may therefore compensate for the temperature drop caused by a local increase in neon density. The fact that most deuterium is deposited near the center could also serve to suppress the runaway generation favored by the presence of neon.

The transported and radiated power throughout a disruption simulation using the injected material densities from both optima, as well as the uniform scenario, is shown in Fig. 5.6b. Note that the difference between the two transport channels is much larger for the spatially varying densities, especially around 10^{-1} ms, which is the timescale of the thermal quench. It is also interesting to note that the Bayesian optimum (dashed) initially has a lower radiated power than the other two, resulting in a higher transported power. This could be due to the fact that a relatively small amount of deuterium is deposited in the outer region which leads to a weaker temperature decrease and, in turn, less effective line radiation from the neon. This difference in initial temperature drop is confirmed when looking at the output data from the respective simulations. The higher transport early on is however compensated for by a lower transport during the majority of the TQ and the early stages of the CQ.

Finally it is once again pertinent to discuss the optimization in a bit more detail. In this increased parameter space, the importance of having a global optimization

Table 5.2: A summary of the parameter values at the different optima obtained via Bayesian optimization and Powell’s method respectively, for different types of disruption simulations. Note that the bar has been removed when denoting the average densities in the profile optimization.

Method	Scenario	$n_D^* (\text{m}^{-3})$	$n_{Ne}^* (\text{m}^{-3})$	c_D^*	c_{Ne}^*
Bayesian	exp. decay	$9.27 \cdot 10^{21}$	$8.51 \cdot 10^{16}$	-	-
	trans. $\sim 0.2\%$	$5.52 \cdot 10^{20}$	$4.30 \cdot 10^{19}$	-	-
	trans. $\sim 0.3\%$	$2.28 \cdot 10^{21}$	$2.04 \cdot 10^{19}$	-	-
	trans. $\sim 0.4\%$	$1.92 \cdot 10^{21}$	$2.78 \cdot 10^{19}$	-	-
	trans. $\sim 0.5\%$	$3.30 \cdot 10^{21}$	$1.32 \cdot 10^{19}$	-	-
	profile	$3.57 \cdot 10^{20}$	$1.08 \cdot 10^{20}$	-6.6	6.3
Powell	exp. decay	$1.17 \cdot 10^{22}$	$7.00 \cdot 10^{16}$	-	-
	trans. $\sim 0.2\%$	$2.79 \cdot 10^{21}$	$1.83 \cdot 10^{19}$	-	-
	trans. $\sim 0.3\%$	$2.36 \cdot 10^{21}$	$1.98 \cdot 10^{19}$	-	-
	trans. $\sim 0.4\%$	$2.61 \cdot 10^{21}$	$1.89 \cdot 10^{19}$	-	-
	trans. $\sim 0.5\%$	$2.71 \cdot 10^{21}$	$2.42 \cdot 10^{19}$	-	-
	profile	$4.06 \cdot 10^{21}$	$5.24 \cdot 10^{19}$	0.44	10

routine became more evident. The minimum found via Powell’s method performs well, but the injected densities given by the two techniques are notably different in nature, especially considering the different signs of c_D . Computationally, it should be mentioned that Powell’s method required a total of 106 simulations to complete. While it converged faster early on, as was expected, the Bayesian sampling made a rather large improvement around the 50th simulation thereby surpassing Powell’s method. Whether it was “luck” is difficult to say based on a single result. Further testing could therefore include running the routine several times with different initial random sampling. One aspect of the Bayesian optimization that so far has been less pronounced compared to previous sections is the construction of a GP, which aside from leading the algorithm towards a minimum, also explores the parameter space and is able to make predictions based on the collected data. In order to visualize the four-dimensional posterior function in terms of marginal distributions we will use a method involving Markov-chain Monte Carlo (MCMC).

By averaging over the different parameters of the constructed posterior mean function, $\mu_n(\bar{n}_D, c_D, \bar{n}_{Ne}, c_{Ne})$, it is possible to illustrate its dependencies with respect to one or two parameters simultaneously. This does however entail numerous multi-dimensional integrals which can be computationally expensive. An efficient numerical technique that is commonly used to compute expectation values within Bayesian statistics is to sample points via the MCMC method. A full description of MCMC falls beyond the scope of this project; for more information the reader is referred to Ref. [51]. In short, the MCMC algorithm uses a *random walk process*, which, given a probability distribution along with a list of starting points, will sample a number of points by taking “steps” from the previously sampled point. From this procedure a discrete distribution of points is obtained, the mean of which is a good estimate of the expectation value, according to *the law of large numbers*;

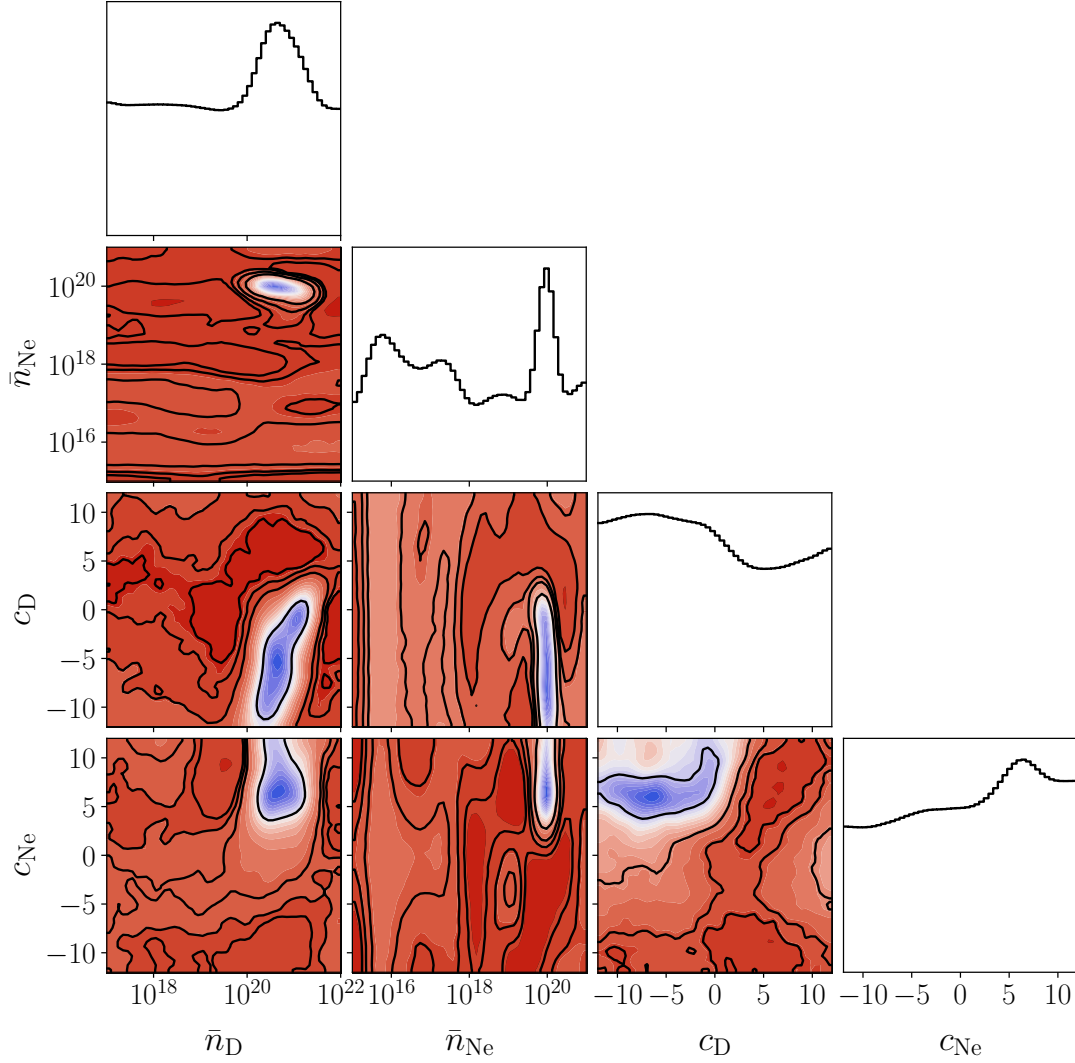


Figure 5.7: Marginal distributions obtained from *Markov chain Monte Carlo* (MCMC) sampling of the posterior probability distribution in Eq. 5.6. Note that the blue contours indicate regions containing more samples (as indicated by the one dimensional histograms on top), implying a smaller value for μ_n .

the accuracy of the method improves with the number of samples. More importantly, the number of samples obtained for specific parameter values will also reveal lower-dimensional marginal distributions. While the posterior mean function is not a probability distribution, it can be subjected to a similar analysis to obtain a distribution that is indicative of the functional shape, from which averages can easily be extracted. The strategy is thus to first sample the parameter space according to a distribution that is inversely proportional to μ_n , and from the obtained set of points determine marginal distributions that can be visualized as one- and two-dimensional histograms. These illustrations will not provide quantitatively accurate values for the objective function, but will rather give a qualitative idea of how the function varies over the studied parameter space.

For this work the MCMC was implemented through the `emcee`¹ python module and the steps were taken based on the *stretch move* [59]. The routine is provided with a function that is proportional to the natural logarithm of the probability distribution up to some additive constant. Here, a heuristic choice was made to use

$$\ln p(\bar{n}_D, c_D, \bar{n}_{Ne}, c_{Ne} | \mathcal{D}_n) \propto -1.8 \cdot \ln \mu_n(\bar{n}_D, c_D, \bar{n}_{Ne}, c_{Ne}), \quad (5.6)$$

since it yielded a good balance during the sampling, where clear preference was given to regions with lower values. The marginal distributions are plotted in a so-called *scatter plot matrix*, using the `corner`² python module; and the result is shown in Fig. 5.7. Note that the color scale in the figure is such that blue regions contain more samples, implying a smaller value for μ_n .

It must be emphasized that the histograms are suitable for qualitative analysis only, for the following reasons. Firstly, the sampled object is the posterior mean function, μ_n , produced by fitting a GP to a data set $\mathcal{D}_n$. This does not take the variance, σ_n^2 , into account, and it is probable that the 525 points contained within the set are not sufficient to adequately explore the large four-dimensional parameter space. Deviations are therefore likely to occur, especially in regions further away from the minimum, since fewer points will have been sampled there during the Bayesian optimization. Secondly, the posterior probability distribution given in Eq. (5.6) is somewhat arbitrary. Specifically, the factor 1.8 weighting the probability, can be increased to exaggerate gradients, or lowered to make the distribution more uniform. Lastly, a Gaussian kernel with $\sigma_{2D} = 1$ and $\sigma_{1D} = 1.5$ is used to artificially smooth out the contour surfaces and one dimensional histograms respectively. The motivation for this is that obtaining naturally smooth distributions would require a very large number of samples combined with high bin counts, which introduces greater computational costs. For these reasons, the number of samples is not included on any axis or colorbar, as a quantitative analysis may be misleading.

Let us now consider the marginal distributions presented in Fig. 5.7. Starting with the average densities, it is interesting to note that they have distinct maxima, with the large peak for $\bar{n}_{Ne}$, being especially sharp. There is however some peculiar behavior in the distribution of the $\bar{n}_{Ne}$, with two peaks located in the region $\bar{n}_{Ne} < 10^{18}$, that previously saw heavy penalization due to incomplete TQs in Sec. 5.2. Furthermore, when studying the two-dimensional distribution for $\bar{n}_{Ne}$ and $\bar{n}_D$, no obvious correlation with the injected deuterium content is found (that is, μ is only weakly varying with $\bar{n}_D$ for a given $\bar{n}_{Ne}$) in this low neon region, which is surprising since a higher deuterium density can assist in achieving a radiative collapse. This is thus likely an example of the previously mentioned deviations that can occur due to the GP being highly uncertain in some regions, leading to inaccurate predictions given by μ_n . This is supported by the fact that there are in total only 48 samples in the parameter space for which $\bar{n}_{Ne} < 10^{18}$, compared to the 477 samples where $\bar{n}_{Ne} \geq 10^{18}$. As such, it is reasonable to exclude these peaks from consideration and note that the preferable average densities seem to lie in the parameter region $\bar{n}_D \sim 10^{20} - 10^{21} \text{ m}^{-3}$ and $\bar{n}_{Ne} \sim 10^{20} \text{ m}^{-3}$.

Moving on to the shaping parameters, these are more smooth in their distribu-

¹emcee: <https://emcee.readthedocs.io/en/stable/>

²corner: <https://corner.readthedocs.io/en/latest/index.html>

tions. For c_{Ne} there is an apparent, positive trend for higher values that agrees with the previous observation that it is beneficial to deposit the majority of neon towards the edge. Even more interesting is the subtle peak around $c_{\text{Ne}} \sim 6$, which is consistent with the value given by the Bayesian optimization, $c_{\text{Ne}}^* = 6.27$. Such a clear preference is not shown in the distribution for c_{D} . This is intriguing when comparing the optima obtained by the two techniques, since they gave vastly different values for this shaping parameter. Furthermore, the location of the optimum given by Powell's, while different from the Bayesian optimum, does in fact lie in the vicinity of the blue regions, and it follows the trends shown in the two-dimensional histograms. For example, comparing the distribution of c_{D} and c_{Ne} , one finds that the blue region protrudes slightly in the top right corner towards higher values of c_{Ne} . This is where the optimum returned by Powell's method is located. In addition, it seems to be preferable to have higher values for $\bar{n}_{\text{D}}$, as well as slightly lower values of n_{Ne} , when c_{D} is increased. This helps explain why the average densities are notably different for Powell's optimum, where $c_{\text{D}}^* = 0.44$.

Figure 5.7 further exemplifies how valuable information, in addition to the location of the optimum, can be extracted from the Bayesian optimization through appropriate use of data processing. In the next section we will see how the routine can be modified to be used for a slightly different application, when studying the sensitivity of the optima obtained in Sec. 5.1 and 5.2.

5.4 Sensitivity analysis

From an engineering perspective, an important quality of any found optimum $\mathbf{x}^*$ is that simulation outputs do not deviate too much from $f(\mathbf{x}^*)$ when perturbing the input parameters by a small amount from the optimum. Allowing for small variations in $\mathbf{x}$ is imperative considering the uncertainties involved in experimental setups, as well as the possibility of the underlying physical model being inaccurate. To address this kind of sensitivity, we instead consider the optimization problem of maximizing f on some small subspace $\mathcal{P}_\varepsilon \subset \mathcal{P}$, centered around $\mathbf{x}^*$, rather than its minimization on the entire $\mathcal{P}$ as previously in this chapter. We therefore refer to this procedure as the *pessimization* of f on $\mathcal{P}_\varepsilon$, which can be thought of simply as the search for the worst case scenario, given $\mathbf{x} \in \mathcal{P}_\varepsilon$. These subspaces are defined such that all points $\mathbf{x} \in \mathcal{P}_\varepsilon$ are within a distance of $\delta x_i := \varepsilon x_i^*$ away from the optimum in each direction, for some fraction $|\varepsilon| < 1$.

The procedure of pessimization is very much like that of the optimizations performed in the previous three sections, but as information on the global behavior of f on $\mathcal{P}_\varepsilon$ is more vital for this assessment, we exclusively use Bayesian optimization in this section and omit the use of Powell's method. Furthermore, with this reason in mind, a larger fraction of the samples are randomly drawn uniformly on this region, compared to previously. At least 50% of all points are sampled this way, while the others are obtained from the expected improvement acquisition function, which guides the sampling of $\mathbf{x}$ towards regions where the GP is either uncertain about f or near the point where it currently believes f attains its maximum.

To illustrate such a procedure, Fig. 5.8 shows the posterior mean function of a GP for the base objective function $f = \mathcal{L}_1$ on a subspace $\mathcal{P}_{10\%}$ centered around

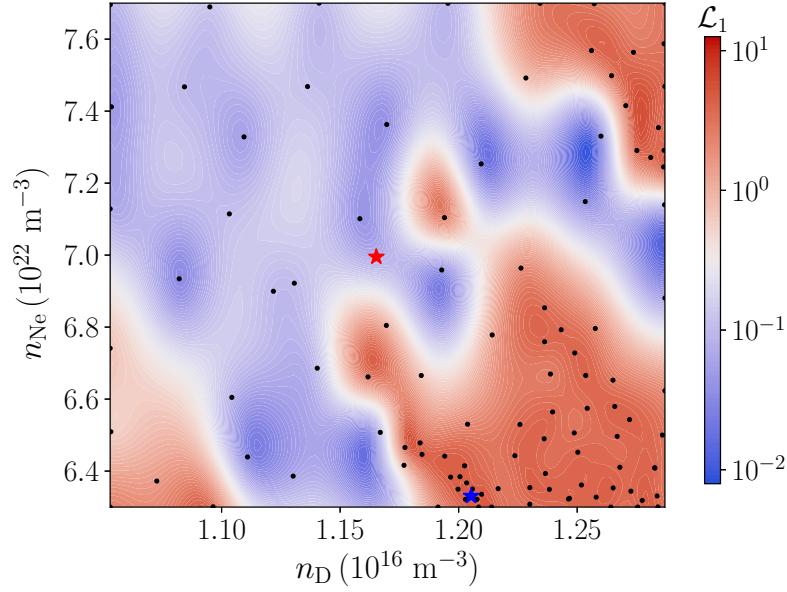


Figure 5.8: Pessimization of the base objective function centered around the optimum (red star) found in Sec. 5.1, from which the input parameters are allowed to vary by 10 %. It illustrates the posterior mean function from a GP based on a total 110 samples (black dots), as well as the current evaluated maximum (blue star), which is $\mathcal{L}_1 \approx 4.5$.

the optimum found in Sec. 5.1 (from the Bayesian optimization). It is based on 60 uniformly distributed samples and an additional 50 points gathered from the acquisition function. From this, it is estimated that the objective function is $\mathcal{L}_1 < 4.5$ for the densities of injected deuterium n_D and neon n_{Ne} within 10 % of the optimum $(n_D^*, n_{Ne}^*) = (1.17 \cdot 10^{16}, 7.00 \cdot 10^{22}) \text{ m}^{-3}$.

There are situations where this way of assessing the sensitivity may perhaps yield a deceptive result, as the assumption of the objective function being continuous might not be true. If the objective function included exceptionally large and localized gradients, or even discontinuities, the GP will apply extremely short correlation lengths in its covariance function. A shorter correlation length implies that a larger fraction of $\mathcal{P}_\varepsilon$ need to be explored to get a reasonable estimate. If such a region is considerably small in size, it can be difficult to locate it, which results in the underestimation of the sensitivity. An example of such a situation is showcased in Fig. 5.9, in which a very narrow peak remains undiscovered for more than 50 function evaluations.

Figure 5.9a shows the posterior mean function based on 50 points that are uniformly distributed on this subspace. Having found that $\mathcal{L}_1 < 0.6$ for all these points, one would at this stage prematurely draw the conclusion that this region is in fact acceptable. As more samples are gathered however, as shown in Fig. 5.9b, the GP finds a confined region close to the border of $\mathcal{P}_{1\%}$ where the objective function rapidly grows towards a value of $\mathcal{L}_1 \approx 3$. After this region is located, the acquisition function guides much of the remaining samples near this peak. Out of all the 100 total samples, the steepest increase between any two points is $\Delta\mathcal{L}_1 = 1.96$ when

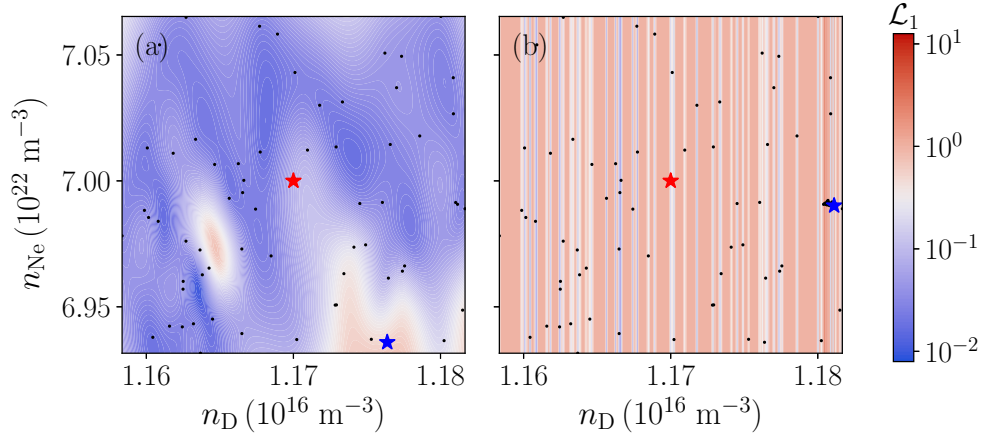


Figure 5.9: Pessimization of the base objective function centered around the minimum (red star) found in Sec. 5.1 using Powell’s method, from which the input parameters are allowed to vary by 1%. The two figures illustrate the posterior mean function from the same GP at two different stages in the optimization procedure, as well as the current evaluated maximum (blue star). Both panels include the same set of 50 uniformly distributed input parameters, from which the evaluated maximum is $\mathcal{L}_1 = 0.53$. An additional 50 points sampled from the expected improvement acquisition function is shown in panel (b), revealing an exceptionally narrow region in which $\mathcal{L}_1 \approx 3$.

$n_D = 1.18112 \cdot 10^{22} \text{ m}^{-3}$ is increased by only 0.004 % and $n_{Ne} = 6.99015 \cdot 10^{16} \text{ m}^{-3}$ by 0.008 %. Having identified this extraordinarily large gradient in the objective function, the correlation length in each direction is dramatically altered in order to account for this behavior. The correlation length in $\log(n_D [\text{m}^{-3}])$ is reduced from $8.42 \cdot 10^{-4}$ to $9.08 \cdot 10^{-6}$ and $\log(n_{Ne} [\text{m}^{-3}])$ is increased from $2.09 \cdot 10^{-3}$ to $1.93 \cdot 10^9$. The posterior mean function is updated and turned into essentially a constant function in n_{Ne} , with rapid oscillations in n_D .

Table 5.3: Output variables corresponding to the minimum on $\mathcal{P}_1$ (maximum on $\mathcal{P}_{10\%}$) of the objective function. The first row is from the disruption scenario considered in Sec. 5.1, where $\mathcal{L}_1$ is used, corresponding to Fig. 5.8. This is followed by the scenarios considered in Sec. 5.2 with four varying magnetic perturbations during the thermal quench, for which the objective $\mathcal{L}_2$ is employed. The corresponding pessimizations of these four cases are included in Appendix B.2.

Scenario	$I_{\text{RE}}^{\text{max}} (\text{A})$	$I_{\Omega}^{\text{final}} (\text{A})$	$\tau_{\text{CQ}} (\text{ms})$	$\eta_{\text{cond}} (\%)$
exp. decay	$1.980 (0.371) \cdot 10^{-3}$	$2.28 (1350) \cdot 10^3$	82.8 (86.4)	-
trans. $\sim 0.2 \%$	$6.55 (5.60) \cdot 10^{-2}$	$2.71 (14.9) \cdot 10^{-6}$	57.3 (58.4)	5.1 (37)
trans. $\sim 0.3 \%$	$1.41 (1.19) \cdot 10^{-2}$	$2.58 (0.92) \cdot 10^{-10}$	55.5 (53.2)	16 (14)
trans. $\sim 0.4 \%$	28.8 (13.5)	$8.94 (1.40) \cdot 10^{-11}$	53.4 (51.1)	26 (22)
trans. $\sim 0.5 \%$	6.26 (4.89)	$2.66 (9.98) \cdot 10^{-12}$	54.6 (52.3)	50 (46)

A summary of the worst performing scenarios found in $\mathcal{P}_{10\%}$ for the respective disruption simulations can be seen in Tab. 5.3, which concludes this chapter. In all cases the obtained output parameters violate the constraints set in Eq. (2.42), by having either a large final Ohmic current or significant transported heat losses. For the prescribed temperature decay, the maximum objective function was located in a region where the final Ohmic current was excessively large, on the order of ~ 1 MA. Furthermore, as was speculated in Sec. 5.2, it seems that a perturbation of the optima for the prescribed heat transport simulations can result in low CQ times or large transported heat losses. It should however be noted that the transported heat losses for most scenarios are lower at the maxima than at their respective optima, while the CQ time is still within the permissible interval. This indicates that the objective function could perhaps be modified, such that it does not return as large values when $\tau_{CQ} > 50$ ms. Finally, we find that $I_{RE}^{\max}$ in all scenarios is still relatively small. However, since the pessimization is conducted with respect to the objective function, this does not exclude the possibility of $I_{RE}^{\max}$ becoming substantial in some region of $\mathcal{P}_{10\%}$.

6

Conclusion and outlook

To achieve reliable tokamak operation, one must navigate a vast parameter space that in some cases extend over several orders of magnitude. Through the use of simulations, it is possible to study different dependencies and create models describing the observed phenomena. Owing to the large number of variables, it is however exceedingly difficult to simultaneously take into account all of the relevant interactions, which often limits the scope of the investigation. By applying an optimization routine that relies on numerics rather than physical reasoning, the process of finding viable parameter regimes for tokamak operation can instead be automated, and made less subjective. Of course, the lack of knowledge about the system requires that the algorithm is guided by suitable constraints and objective functions; but when used appropriately, it can perform a comprehensive search of the parameter space.

In this work we implemented two black-box optimization routines that were subsequently applied to simulations of disruption scenarios representative of ITER. Using objective functions that consider the maximum runaway current, the remaining Ohmic current, the current quench time, and the transported heat losses, we are able to find preferable parameter values for the amount of injected deuterium and neon during a MMI. In this chapter we will provide a summary of the results, including the obtained optimal parameters and the performance of the two optimization methods, as well as give an outlook for possible improvements and further studies.

6.1 Conclusion

As a proof of concept, we first optimized a disruption scenario where the temperature had a prescribed exponential decay during the thermal quench and the injected density profiles were uniform. In this case both methods converged towards similar minima, around $n_D \sim 10^{22} \text{ m}^{-3}$ and $n_{Ne} \sim 8 \cdot 10^{16} \text{ m}^{-3}$. For this, an objective function was constructed which depended on the maximal runaway current, remaining Ohmic current, and current quench time. All these parameters were found to be within the acceptable limits simultaneously at the optima found by the two optimization methods. We find that in the optimal cases the relatively large amount of injected deuterium resulted in an efficient suppression of REs, due to the increased drag force, while also contributing to the temperature decrease such that an adequate current quench time below 150 ms could be obtained. A sensitivity analysis in the form of pessimization was also conducted, where the objective function was

maximized in a region where the parameters assumed values within 10 % of the Bayesian minimum. This exercise revealed that the final Ohmic current varied greatly within this subspace, reaching ~ 1 MA at the maximum. Furthermore, when a similar analysis was done for a contracted region corresponding to 1 % variation of parameter values, a sharp peak was found, indicating discontinuities in the function, related to the formation of Ohmic current channels in an unpredictable fashion.

The relatively low amount of injected neon at the respective minima could be explained by the fact that the utilized objective function did not take transported heat losses into account. To relax this limitation, the disruption simulations were modified to use a self-consistent temperature evolution during the thermal quench, from which the transported fraction of the initial thermal energy could be calculated. With an updated objective function, the same optimization routines were applied to scenarios where heat and runaway particle transport was prescribed consistently, through setting magnetic perturbation strengths in Rechester-Rosenbluth type diffusion models. This complicated the dynamics of the thermal quench, leading to incomplete thermal quenches across a large region of the parameter space. It was noted that the final Ohmic current was vanishingly small at the resulting minima. Furthermore, the optimal neon densities were around three orders of magnitude greater than that of the previous minima. Since the maximal runaway current was sufficiently low for all cases, it was surmised that the diffusion of runaway electrons was sufficient to limit the runaway seed, and thus significantly lower maximum runaway currents, despite the increased amount of injected neon. However, the respective current quench times were in the vicinity of the lower limit, and the transported heat losses became substantial for higher perturbation strengths. A sensitivity analysis indicated that the latter quantities could assume even less agreeable values within a 10 % variation of the optima.

Hoping to improve the worst performing scenario with prescribed transport, for which $\delta B/B = 0.5\%$, two additional degrees of freedom were included in the optimization, corresponding to parameters governing the profile of the injected material. This improved performance greatly, with both optima returning satisfactory values for the output parameters. Most notably, the transported fraction decreased substantially, likely owing to the fact that the average neon density at the optima was higher than the previously obtained values. Furthermore, the optimal profiles were such that the majority of neon was deposited near the plasma edge. Since runaway generation is generally more prominent near the center of the plasma, it is possible that an increase in the amount of neon does not induce excessive RE currents, when distributed in this manner.

Regarding the two optimization schemes, we may conclude that they have different strengths and weaknesses, which can be considered when choosing an optimization method. Powell's method appears to converge faster towards a minimum in most cases, though not during the four-dimensional optimization. A proper comparison between the efficiency of the methods would require further examination, considering for example different starting positions for Powell's method and different random states for the Bayesian method. It did however become evident, specifically for the profile optimization, that Powell's method could struggle with getting stuck in local optima. The Bayesian approach did not encounter the same problem. In

addition, the Gaussian processes obtained during the procedure were highly informative and provided a great basis for further insight.

6.2 Outlook

One aspect of this project that required much consideration was how to limit the scope, since there are numerous ideas for improvements and further research. To start with, the respective optimization routines can be modified in several ways in order to increase efficiency and prevent complications. With regards to Powell's method, there are techniques not used here that can be implemented to avoid linear dependence, that makes better use of the constructed basis set [60]. In addition, methods such as *simulated annealing* can be applied to avoid problems with the routine getting stuck in local minima [51].

As for the Bayesian optimization, it is always possible to further improve the prior distribution function, used to encapsulate any possible prior knowledge of the objective function and integrate it into the Gaussian process. Naturally, how this is to be improved upon, heavily depends on the situation, such as the black-box model or the physical scenario considered. In this work we placed the standard normal distribution as prior for the logarithm of the objective function at each point in parameter space to force it to yield positive real numbers. One could perhaps attempt to incorporate the fact that for some input parameters, the objective function tends to increase when moving towards the boundary in parameter space, assuming it has been constructed such that the optimum is situated well inside the boundary. This should be able to increase the efficiency of the routine in some cases.

Another possible addition would be to analyze the physical model used in this work. The simulations were conducted using a fully fluid model of the plasma, which only resolved one spatial dimension. One simple improvement would be to assess the obtained optima through similar simulations that instead use the fully kinetic mode in DREAM, which also resolves the electron momentum space [41]. It should however be noted that there are several additional levels of approximation in between, such as a fluid model that also takes the momentum dependence of runaway transport into account [61].

The numerical experiments can also be developed further. As was pointed out, the self-consistent evolution of magnetic perturbations is a complex phenomenon and the choices made here to model them are, to a large degree, arbitrary. Coupling the simulations to an MHD code that evolves the magnetic field self-consistently would increase the computational expense of the simulations beyond what is reasonable for optimization. By studying the perturbations with frameworks such as JOREK, it might however be possible to make an informed choice, modelling and prescribing the transport more appropriately [58]. Furthermore, there are properties of the material injection that were not considered in this work. In the simulations it was assumed that all of the injected material was instantaneously deposited according to a certain profile at time $t = 0$. In practice the material is however injected into the plasma through methods such as *massive gas injection* or *pellet injection*, which introduces new dynamics to the system [62]. In addition, it is common to use a two-stage injection scheme, where the deuterium and neon would be inserted at different

points in time during the disruption. This entails new degrees of freedom which are suitable for optimization.

Finally, in this project we have focused on a material injection scenario, but there are many other properties of tokamak operation that could be subjected to a similar analysis. It is for example possible that the initial magnetic geometry could have a significant impact during disruption events. Similarly, the initial temperature and current profiles are integral to the nature of a disruption (the avalanche mechanism is for example exponentially sensitive to the total Ohmic current), and could therefore also be targets of optimization or sensitivity analysis. Moreover, it is important to note that the optimization methods are not exclusive to disruption scenarios. While the results obtained in this project shows promise for future disruption mitigation in tokamak reactors, they have more importantly illustrated the feasibility of an optimization approach to this type of problem. One could for example imagine using the same framework to study tokamak start-up or steady-state operation, but the possibilities extend beyond this as well [63]. The results in this work indicate that through proper use of optimization techniques it is possible to make the black-box, along with the intricacies of fusion energy, more transparent.

Bibliography

- [1] Intergovernmental Panel on Climate Change (IPCC): “Climate change 2021: The physical science basis.” (2021). URL: <https://www.ipcc.ch/report/ar6/wg1/>.
- [2] World Meteorological Organization (WMO): “State of the global climate 2021” (2022). URL: <https://wedocs.unep.org/20.500.11822/40033>.
- [3] K. Forinash: *Physics and the Environment*. Morgan & Claypool Publishers (2017). ISBN: 978-1-6817-4493-3. DOI: 10.1088/978-1-6817-4493-3.
- [4] J. Freidberg: *Plasma Physics and Fusion Energy*. Cambridge University Press (2008). ISBN: 9781139462150.
- [5] Merriam-Webster.com Dictionary: “Tokamak”. Accessed 20 May 2022, URL: <https://www.merriam-webster.com/dictionary/tokamak>.
- [6] F. Chen: *Introduction to Plasma Physics and Controlled Fusion*. Springer International Publishing (1974).
- [7] E. M. Hollmann, P. B. Aleynikov, T. Fülöp, D. A. Humphreys, V. A. Izzo, M. Lehen, V. E. Lukash, G. Papp, G. Pautasso, F. Saint-Laurent & J. A. Snipes: “Status of research toward the ITER disruption mitigation system”. *Physics of Plasmas*, **22**, 021802 (2015). DOI: 10.1063/1.4901251.
- [8] S. Li, H. Jiang, Z. Ren & C. Xu: “Optimal tracking for a divergent-type parabolic PDE system in current profile control”. *Abstract and Applied Analysis* (2014). DOI: 10.1155/2014/940965.
- [9] T. Boyd & J. Sanderson: *The Physics of Plasmas*. Cambridge University Press (2003). DOI: 10.1017/CB09780511755750.
- [10] D. Gurnett & A. Bhattacharjee: *Introduction to Plasma Physics: With Space and Laboratory Applications*. Cambridge University Press (2005).
- [11] A. Gurevich, G. Milikh & R. Roussel-Dupre: “Runaway electron mechanism of air breakdown and preconditioning during a thunderstorm”. *Physics Letters A*, **165**, 463–468 (1992). ISSN 0375-9601. DOI: 10.1016/0375-9601(92)90348-P.
- [12] C. Kennel: “Consequences of a magnetospheric plasma”. *Reviews of Geophysics*, **7**, 379–419 (1969). DOI: 10.1029/RG007i001p00379.
- [13] A. Morozov: *Introduction to Plasma Dynamics*. Taylor & Francis (2012). ISBN: 9781439881323.
- [14] P. Helander & D. Sigmar: *Collisional Transport in Magnetized Plasmas*. Cambridge Monographs on Plasma Physics. Cambridge University Press (2005). ISBN: 9780521020985.
- [15] H. Alfvén: “Existence of electromagnetic-hydrodynamic waves”. *Nature*, **150**, 405–406 (1942). DOI: 10.1038/150405d0.

- [16] H. Dreicer: “Electron and ion runaway in a fully ionized gas. I”. *Phys. Rev.*, **115**, 238–249 (1959). DOI: 10.1103/PhysRev.115.238.
- [17] H. Dreicer: “Electron and ion runaway in a fully ionized gas. II”. *Phys. Rev.*, **117**, 329–342 (1960). DOI: 10.1103/PhysRev.117.329.
- [18] J. Connor & R. Hastie: “Relativistic limitations on runaway electrons”. *Nuclear Fusion*, **15**, 415–424 (1975). DOI: 10.1088/0029-5515/15/3/007.
- [19] J. R. Martín-Solís, R. Sánchez & B. Esposito: “Experimental observation of increased threshold electric field for runaway generation due to synchrotron radiation losses in the FTU tokamak”. *Phys. Rev. Lett.*, **105**, 185002 (2010). DOI: 10.1103/PhysRevLett.105.185002.
- [20] A. Stahl, E. Hirvijoki, J. Decker, O. Embréus & T. Fülöp: “Effective critical electric field for runaway-electron generation”. *Phys. Rev. Lett.*, **114**, 115002 (2015). DOI: 10.1103/PhysRevLett.114.115002.
- [21] L. Hesslow, O. Embréus, G. J. Wilkie, G. Papp & T. Fülöp: “Effect of partially ionized impurities and radiation on the effective critical electric field for runaway generation”. *Plasma Physics and Controlled Fusion*, **60**, 074010 (2018). DOI: 10.1088/1361-6587/aac33e.
- [22] L. Hesslow, O. Embréus, M. Hoppe, T. C. DuBois, G. Papp, M. Rahm & T. Fülöp: “Generalized collision operator for fast electrons interacting with partially ionized impurities”. *Journal of Plasma Physics*, **84** (2018). DOI: 10.1017/s0022377818001113.
- [23] L. Hesslow, O. Embréus, O. Vallhagen & T. Fülöp: “Influence of massive material injection on avalanche runaway generation during tokamak disruptions”. *Nuclear Fusion*, **59**, 084004 (2019). DOI: 10.1088/1741-4326/ab26c2.
- [24] J. Martín-Solís, A. Loarte & M. Lehnen: “Formation and termination of runaway beams in ITER disruptions”. *Nuclear Fusion*, **57**, 066025 (2017). DOI: 10.1088/1741-4326/aa6939.
- [25] O. Vallhagen, O. Embreus, I. Pusztai, L. Hesslow & T. Fülöp: “Runaway dynamics in the DT phase of ITER operations in the presence of massive material injection”. *Journal of Plasma Physics*, **86**, 475860401 (2020). DOI: 10.1017/S0022377820000859.
- [26] T. Fülöp, P. Helander, O. Vallhagen, O. Embreus, L. Hesslow, P. Svensson, A. Creely, N. Howard & P. Rodriguez-Fernandez: “Effect of plasma elongation on current dynamics during tokamak disruptions”. *Journal of Plasma Physics*, **86**, 474860101 (2020). DOI: 10.1017/S002237782000001X.
- [27] P. Helander, H. Smith, T. Fülöp & L.-G. Eriksson: “Electron kinetics in a cooling plasma”. *Physics of Plasmas*, **11**, 5704–5709 (2004). DOI: 10.1063/1.1812759.
- [28] H. Smith, P. Helander, L.-G. Eriksson & T. Fülöp: “Runaway electron generation in a cooling plasma”. *Physics of Plasmas*, **12**, 122505 (2005). DOI: 10.1063/1.2148966.
- [29] P. Aleynikov & B. N. Breizman: “Generation of runaway electrons during the thermal quench in tokamaks”. *Nuclear Fusion*, **57**, 046009 (2017). DOI: 10.1088/1741-4326/aa5895.
- [30] I. Svenningsson, O. Embreus, M. Hoppe, S. L. Newton & T. Fülöp: “Hot-tail

- runaway seed landscape during the thermal quench in tokamaks”. *Physical Review Letters*, **127**, 035001 (2021). DOI: 10.1103/PhysRevLett.127.035001.
- [31] M. Rosenbluth & S. Putvinski: “Theory for avalanche of runaway electrons in tokamaks”. *Nuclear Fusion*, **37**, 1355–1362 (1997). DOI: 10.1088/0029-5515/37/10/i03.
- [32] H. Grad & H. Rubin: “Hydromagnetic equilibria and force-free fields”. *Proceedings of the 2nd UN Conference on the Peaceful Uses of Atomic Energy*, **31**, 190 (1958).
- [33] V. D. Shafranov: “Plasma equilibrium in a magnetic field”. *Reviews of Plasma Physics*, **2**, 103 (1966).
- [34] J. Wesson: *Tokamaks; 4th edition*. International series of monographs on physics. Oxford Univ. Press, Oxford (2011).
- [35] B. N. Breizman: “Marginal stability model for the decay of runaway electron current”. *Nuclear Fusion*, **54**, 072002 (2014). DOI: 10.1088/0029-5515/54/7/072002.
- [36] M. Hoppe: *Runaway-electron model development and validation in tokamaks*. Phd thesis, Chalmers University of Technology (2021). URL: <https://research.chalmers.se/publication/527630>.
- [37] I. Pusztai, M. Hoppe & O. Vallhagen: “Runaway dynamics in disruptions with current relaxation” (2022). *Submitted to the Journal of Plasma Physics*, arXiv: 2206.00904.
- [38] M. Lehnen & the ITER Disruption Mitigation Task Force: “The ITER disruption mitigation system – design progress and design validation” (2021). Presented during Theory & Simulations of Disruptions Workshop, PPPL, URL: <https://tsdw.pppl.gov/Talks/2021/Lehnen.pdf>.
- [39] R. Tinguely, V. Izzo, D. Garnier, A. Sundström, K. Särkimäki, O. Embréus, T. Fülöp, R. Granetz, M. Hoppe, I. Pusztai & R. Sweeney: “Modeling the complete prevention of disruption-generated runaway electron beam formation with a passive 3D coil in SPARC”. *Nuclear Fusion*, **61**, 124003 (2021). DOI: 10.1088/1741-4326/ac31d7.
- [40] G. Papp, M. Drevlak, T. Fülöp, P. Helander & G. I. Pokol: “Runaway electron losses caused by resonant magnetic perturbations in ITER”. *Plasma Physics and Controlled Fusion*, **53**, 095004 (2011). DOI: 10.1088/0741-3335/53/9/095004.
- [41] M. Hoppe, O. Embréus & T. Fülöp: “DREAM: A fluid-kinetic framework for tokamak disruption runaway electron simulations”. *Computer Physics Communications*, **268**, 108098 (2021). DOI: 10.1016/j.cpc.2021.108098.
- [42] A. Redl, C. Angioni, E. Belli & O. Sauter: “A new set of analytical formulae for the computation of the bootstrap current and the neoclassical conductivity in tokamaks”. *Physics of Plasmas*, **28**, 022502 (2021). DOI: 10.1063/5.0012664.
- [43] H. P. Summer: “The ADAS user manual, version 2.6” (2004). URL: <http://www.adas.ac.uk>.
- [44] D. Reiter: “The data file AMJUEL: Additional atomic and molecular data for EIRENE” (2000). URL: <http://www.eirene.de>.
- [45] A. B. Rechester & M. N. Rosenbluth: “Electron heat transport in a tokamak with destroyed magnetic surfaces”. *Phys. Rev. Lett.*, **40**, 38–41 (1978). DOI: 10.1103/PhysRevLett.40.38.

- [46] K. Särkimäki, O. Embreus, E. Nardon, T. Fülöp & J. contributors: “Assessing energy dependence of the transport of relativistic electrons in perturbed magnetic fields with orbit-following simulations”. *Nuclear Fusion*, **60**, 126050 (2020). DOI: 10.1088/1741-4326/abb9e9.
- [47] L. Hesslow, L. Unnerfelt, O. Vallhagen, O. Embreus, M. Hoppe, G. Papp & T. Fülöp: “Evaluation of the Dreicer runaway generation rate in the presence of high- Z impurities using a neural network”. *Journal of Plasma Physics*, **85**, 475850601 (2019). DOI: 10.1017/S0022377819000874.
- [48] I. Svenningsson: *Hot-tail runaway electron generation in cooling fusion plasmas*. Master’s thesis, Chalmers University of Technology (2020). DOI: 20.500.12380/300899.
- [49] H. M. Smith & E. Verwichte: “Hot tail runaway electron generation in tokamak disruptions”. *Physics of Plasmas*, **15**, 072502 (2008). DOI: 10.1063/1.2949692.
- [50] R. J. LeVeque *et al.*: *Finite volume methods for hyperbolic problems*, vol. 31. Cambridge university press (2002).
- [51] W. Press, S. Teukolsky, W. Vetterling & B. Flannery: *Numerical Recipes 3rd Edition: The Art of Scientific Computing*. Cambridge University Press (2007). ISBN: 9780521880688.
- [52] P. I. Frazier: “A tutorial on Bayesian optimization” (2018). DOI: 10.48550/ARXIV.1807.02811.
- [53] M. J. D. Powell: “An efficient method for finding the minimum of a function of several variables without calculating derivatives”. *The Computer Journal*, **7**, 155–162 (1964). ISSN 0010-4620. DOI: 10.1093/comjnl/7.2.155.
- [54] C. E. Rasmussen & C. K. I. Williams: *Gaussian processes for machine learning*. Adaptive computation and machine learning. MIT Press (2006). ISBN: 026218253X.
- [55] D. R. Jones, M. Schonlau & W. J. Welch: “Efficient global optimization of expensive black-box functions”. *Journal of Global optimization*, **13**, 455–492 (1998). DOI: 10.1023/A:1008306431147.
- [56] S. Putvinski, N. Fujisawa, D. Post, N. Putvinskaya, M. Rosenbluth & J. Wesley: “Impurity fueling to terminate tokamak discharges”. *Journal of Nuclear Materials*, **241–243**, 316–321 (1997). ISSN 0022-3115. DOI: 10.1016/S0022-3115(97)80056-6.
- [57] T. Fehér, H. M. Smith, T. Fülöp & K. Gál: “Simulation of runaway electron generation during plasma shutdown by impurity injection in ITER”. *Plasma Physics and Controlled Fusion*, **53**, 035014 (2011). DOI: 10.1088/0741-3335/53/3/035014.
- [58] M. Hoelzl, G. Huijsmans, S. Pamela, M. Bécoulet, E. Nardon, F. Artola, B. Nkonga, C. Atanasiu, V. Bandaru, A. Bhole, D. Bonfiglio, A. Cathey, O. Czarny, A. Dvornova, T. Fehér, A. Fil, E. Franck, S. Futatani, M. Gruca, H. Guillard, J. Haverkort, I. Holod, D. Hu, S. Kim, S. Korving, L. Kos, I. Krebs, L. Kripner, G. Latu, F. Liu, P. Merkel, D. Meshcheriakov, V. Mitterauer, S. Mochalsky, J. Morales, R. Nies, N. Nikulsin, F. Orain, J. Pratt, R. Ramasamy, P. Ramet, C. Reux, K. Särkimäki, N. Schwarz, P. S. Verma, S. Smith, C. Sommariva, E. Strumberger, D. van Vugt, M. Verbeek, E. Westerhof, F. Wieschollek & J. Zielinski: “The JOREK non-linear extended

- MHD code and applications to large-scale instabilities and their control in magnetically confined fusion plasmas”. *Nuclear Fusion*, **61**, 065001 (2021). DOI: 10.1088/1741-4326/abf99f.
- [59] J. Goodman & J. Weare: “Ensemble samplers with affine invariance”. *Communications in Applied Mathematics and Computational Science*, **5**, 65 – 80 (2010). DOI: 10.2140/camcos.2010.5.65.
- [60] R. Brent: *Algorithms for Minimization Without Derivatives*. Prentice-Hall series in automatic computation. Prentice-Hall (1972). ISBN: 9780130223357.
- [61] P. Svensson, O. Embreus, S. L. Newton, K. Särkimäki, O. Vallhagen & T. Fülöp: “Effects of magnetic perturbations and radiation on the runaway avalanche”. *Journal of Plasma Physics*, **87**, 905870207 (2021). DOI: 10.1017/S0022377820001592.
- [62] O. Vallhagen, I. Pusztai, M. Hoppe, S. L. Newton & T. Fülöp: “Effect of two-stage shattered pellet injection on tokamak disruptions”. *Nuclear Fusion* (2022). DOI: 10.48550/ARXIV.2201.10279.
- [63] M. Hoppe, I. Ekmark, E. Berger & T. Fülöp: “Runaway electron generation during tokamak start-up”. *Journal of Plasma Physics* (2022). DOI: 10.48550/ARXIV.2203.09900.

A

Analytic magnetic geometry

In DREAM it is assumed that every particle trajectory exactly follows magnetic field lines, lying in surfaces of constant poloidal flux $\psi(r)$. The radial coordinate r is defined as the half width of a flux surface in the midplane, and is used together with the two angles θ, φ to describe any position $\mathbf{x} = \mathbf{x}(r, \theta, \varphi)$ in this toroidal geometry. As shown in Fig. A.1, θ and φ are angles for the poloidal and toroidal directions, respectively.

More specifically we consider an up-down symmetric magnetic geometry. Each flux surface is parameterized by a major radius R_0 , a reference poloidal flux $\psi_{\text{ref}}(r)$ as well as the three shaping parameters $\kappa(r)$ *elongation*, $\delta(r)$ *triangularity* and $\Delta(r)$ *Shafranov shift*, as provided in Tab. A.1. Firstly, the elongation is a scale factor determining the ellipticity of a flux surface and is usually $\kappa \gtrsim 1$. The triangularity $\delta(r)$ determines the horizontal shift of the upper and lowermost points in the z -direction and finally, the Shafranov shift $|\Delta(r)| < a$, is a shift of the flux surfaces in the direction of the major radius. Using the Cartesian basis vectors, any point $\mathbf{x}$ is specified by

$$\begin{aligned}\mathbf{x} &= R(r, \theta) \hat{\mathbf{R}}(\varphi) + z(r, \theta) \hat{\mathbf{z}}, \\ R(r, \theta) &= R_0 + \Delta(r) + r \cos[\theta + \delta(r) \sin \theta], \\ z(r, \theta) &= r \kappa(r) \sin \theta, \\ \hat{\mathbf{R}}(\varphi) &= \cos \varphi \hat{\mathbf{x}} + \sin \varphi \hat{\mathbf{y}}.\end{aligned}\tag{A.1}$$

Note that this parameterization makes the angle θ slightly differ from the angle in the poloidal plane between the outer midplane and any position $\mathbf{x} = \mathbf{x}(r, \theta, \varphi)$, with the origin at the magnetic axis. The Jacobian for this coordinate system is given by

$$\mathcal{J} = \left| \frac{\partial \mathbf{x}}{\partial r} \cdot \left(\frac{\partial \mathbf{x}}{\partial \theta} \times \frac{\partial \mathbf{x}}{\partial \varphi} \right) \right| = \frac{1}{|\nabla r \cdot (\nabla \theta \times \nabla \varphi)|},\tag{A.2}$$

which can be evaluated in closed form (see appendix of Ref. [41]) in terms of the parameterization in Eq. (A.1). The area V' introduced in Ch. 3 is the Jacobian of the reduced one-dimensional coordinate system of only r , which is used when calculating the flux surface average $\langle X \rangle$ of quantity X , and is given by

$$\langle X \rangle := \frac{1}{V'} \iint \mathcal{J} d\theta d\varphi X(r, \theta, \varphi),\tag{A.3}$$

$$V' := \iint \mathcal{J} d\theta d\varphi.\tag{A.4}$$

The integrals are evaluated from 0 to 2π in both angles θ and φ .

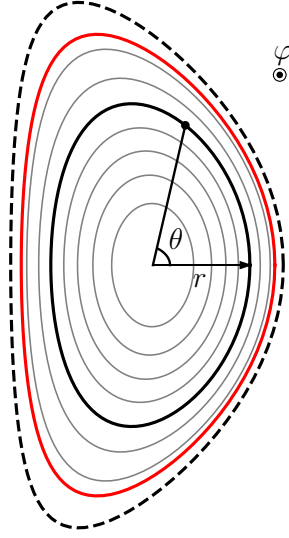


Figure A.1: Illustration of the three coordinates used to describe positions in a shaped tokamak magnetic geometry, namely the radial coordinate r , poloidal angle θ and the toroidal angle φ (out of the paper). This particular geometry uses the shaping profiles in Tab. A.1, characterizing that of a ITER-like plasma, which is considered in Ch. 5. The outermost flux surface (dashed black) indicates the tokamak wall at $r = 2.15$ m. Inside of this is the edge of the plasma at $r = 2$ m (solid red) as well as a surface inside the plasma labeled with r (solid black).

Being symmetric in the toroidal direction, i.e. any derivatives $\partial/\partial\varphi$ vanish, the magnetic field can be expressed in the following form

$$\mathbf{B} = G(r)\nabla\varphi + \frac{\psi'_{\text{ref}}(r)}{2\pi}\nabla\varphi \times \nabla r, \quad (\text{A.5})$$

where the function $G(r)$ describes the toroidal component of $\mathbf{B}$ and ψ'_{ref} is the derivative of the reference poloidal flux ψ_{ref} . Note that the actual poloidal flux ψ varies in a simulation, as the current density changes in time, recall Eq. (3.3), but the magnetic geometry, including the reference poloidal flux is held fixed in the simulation. The magnetic field strength B is assumed to have exactly one maximum and one minimum, per flux surface. Usually in tokamaks, as is the case for the considered analytic magnetic geometry of ITER described in Tab. A.1, the toroidal component of the magnetic field is much larger than the poloidal component. Therefore $B \approx |G(r)\nabla\varphi| = G(r)/R \implies B \propto 1/R$ approximately for a slowly varying $G(r)$.

Table A.1: Polynomial coefficients describing the plasma and its shaping profiles. Together with the following coefficients, each radially varying quantity is obtained by adding up the terms $a_n r^n$ over integers n . These sets the magnetic geometry similar to that of ITER [37]. The major radius is set to $R_0 = 6$ m, and the outermost minor radius is $a = 2$ m.

Quantity	a_0	a_1	a_2	a_3	a_4
ψ_{ref}/R_0 [Tm]	-	-	0.794102	-0.117139	-
κ	1.5	-	-	-	0.02
δ	-	-	0.035	-	0.017
Δ [m]	-	-	-0.00658678	-0.00634124	-
G/R_0 [T]	5.3	-	-0.310741	0.107719	-

B

Gaussian process predictions

With the use of GP regression, which is detailed in Sec. 4.2.1, it is possible to obtain a point estimate — the posterior mean function μ_n — estimate of the objective function in question, provided a set of n function evaluations. In Sec. 5.1, this method was applied to the function $\mathcal{L}_1$, as defined in Eq. (5.2), being a function of the maximum RE current, final Ohmic current and the CQ time. All of these are outputs of fluid simulations of the disruption scenarios, as described in Ch. 3, which in turn depend non-linearly on a set of input parameters describing various initial conditions and other aspects of a disruption simulation. All but two inputs were set according to a baseline scenario. These two, namely those of injected neon and deuterium densities n_{Ne} and n_{D} , respectively, were instead treated as variables when optimizing $\mathcal{L}_1$.

Included here are two aspects of the GP regressions performed in Ch. 5, namely the evolution of the posterior mean function, towards in its final form in Fig. 5.2b in Sec. 5.1, followed by the four sensitivity analyses presented in Tab. 5.3. The former is to illustrate an example where the Bayesian optimization method first finds a local optimum and is then able to locate a region with an even lower value in $\mathcal{L}_1$.

B.1 Initially prescribed temperature evolution

A total of 200 samples were used to construct the GP illustrated in Fig. 5.2. These were chosen based on the acquisition function described in Sec. 4.2.2, weighing between exploration of the parameter space and minimization of the objective function. To illustrate the sampling procedure and the evolution of the GP, Fig. B.1 shows the posterior mean function at different stages during the Bayesian optimization in Sec. 5.4. Notably, after 100 samples it is evident that the routine has found a preferred region, corresponding to rather high amounts of injected deuterium and neon. This is where most sampled points lie. Large parts of the parameter space are however still sparsely sampled at this point and it can be seen that after 115 samples, a new minimum is found in a region with significantly lower neon densities. What follows is an extensive sampling of the newly found region where the posterior mean function is low and the uncertainty is high. After 130 samples the posterior mean function has become more nuanced within this small part of the parameter space, and after the full 200 samples a narrow blue region has appeared, where the GP predicts that the objective function is on the order of $\sim 10^{-1}$.

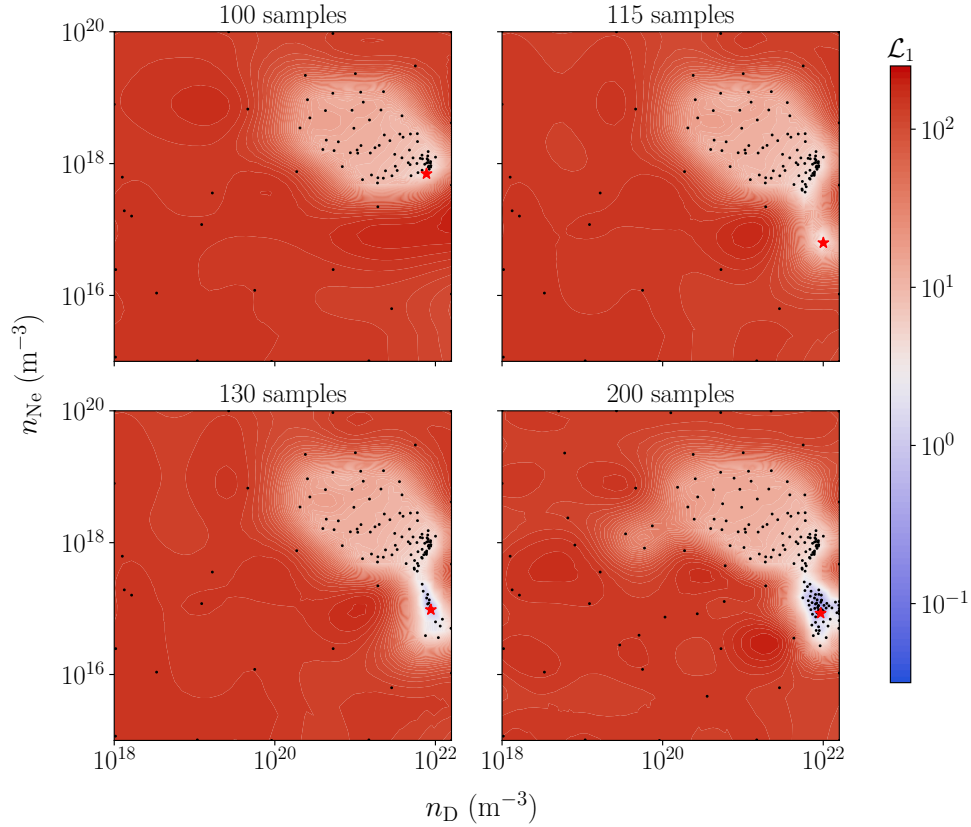


Figure B.1: GPs with increasing number of samples of the objective function $\mathcal{L}_1$ included. The red star shows the sample with the current evaluated optimum $\mathbf{x}_n^*$, and the black dots indicate each of the n samples.

B.2 Pessimizations of found minima

A sensitivity analysis, analogous to that in Sec. 5.4, was conducted for all scenarios with transported heat losses and uniform material injection. The posterior mean function of a GP constructed by sampling the corresponding $\mathcal{P}_{10\%}$ subspaces is shown in Fig. B.2. It can be seen that, in the vicinity of the optimum obtained for $\delta B/B = 0.3\%$ and $\delta B/B = 0.5\%$, the objective function, $\mathcal{L}_2$, increases significantly for higher amounts of injected deuterium and neon. This is likely a result of the disruptions having a lower CQ time when injecting more material. A similar decrease of the CQ time was obtained at the maximum for $\delta B/B = 0.4\%$, which, in contrast to the previous cases, is located in a region of lower deuterium density. It is possible that this is a consequence of the non-linear behavior of the radiated power. As for the optimum obtained when $\delta B/B = 0.2\%$ (which, as was seen in Fig. 5.4, is located in a slightly different region than the other optima) the maximum value of the objective function found in $\mathcal{P}_{10\%}$ instead corresponded to lower values of n_D and n_{Ne} . From what is shown in Tab. 5.3, it seems that this is a result of the transported heat losses increasing, rather than a decrease in the CQ time.

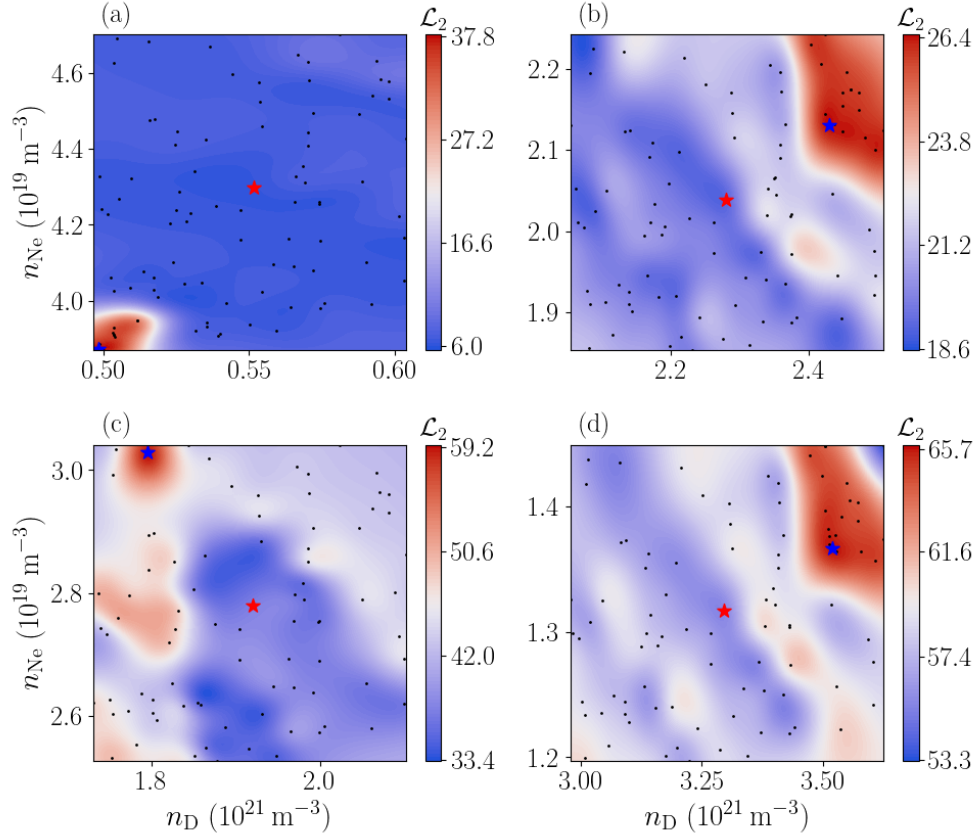


Figure B.2: Pessimizations of the base objective function centered around the optimum (red star) found in Sec. 5.2, from which the input parameters are allowed to vary by 10 %. The subfigures correspond to magnetic perturbations $\delta B/B$ of (a) 0.2 %, (b) 0.3 %, (c) 0.4 % and (d) 0.5 %, all radially uniform. The output variables of maximum RE current, final Ohmmic current, CQ time and the transported fraction for the optimum, as well as for the maximizer (blue star) on $\mathcal{P}_{10\%}$, are included in Tab. 5.3.

DEPARTMENT OF PHYSICS
CHALMERS UNIVERSITY OF TECHNOLOGY
Gothenburg, Sweden
www.chalmers.se



CHALMERS
UNIVERSITY OF TECHNOLOGY