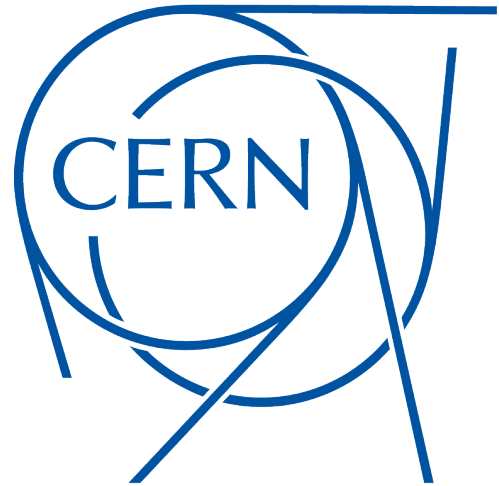


POLITECNICO DI TORINO

Master degree in "Mathematical engineering"

Numerical Methods to Optimize the LHC Luminosity Computation



Supervisors

Stefano Berrone
Guido Sterbini

Student

Matteo Rufolo - 288520

Academic year 2023-2024

Abstract

One of the most significant merit figures in the Large Hadron Collider, LHC, is the luminosity. Larger values of luminosity allow for more interesting physical observations by the experiments, and it is a metric to quantify the instantaneous rate of the particles interacting in the collider. The luminosity is a scalar observable that indicates how much “Physics” is created during the run in the LHC, hence it is directly connected to LHC’s capability for discovery. The intensity of the two beams, the number of colliding bunches, the bunch profiles, the frequency of revolution, etc. . . are all physical quantities that contribute to the luminosity. Several facets of the accelerator and the beam properties are integrated into this figure.

The objective of this thesis is to employ various numerical and mathematical techniques in order to enhance the efficiency of luminosity computation. At the beginning, we will present a mathematical model for the calculation of colliding bunches, followed by a comprehensive examination of an optimized Python implementation. Subsequently, our objective will be to optimize the number of colliding bunches through the strategic arrangement of these bunches, more focusing on a number theory approach. The primary difficulty is in the vast number of potential combinations and the numerous limits that must be adhered to. Consequently, achieving our fundamental goal becomes almost unreachable, as we shall soon discover. Subsequently, we will analyse the bunch profiles, attempting to calculate them based on the observed luminosity. The latter refers to a scalar quantity that is derived through a convolution in three-dimensional space and time of these bunch profiles. Our objective is to obtain distinct values from a single scalar using an analytical approach. This presents a problem that is inherently ill-conditioned. However, as we will demonstrate later, we are able to obtain intriguing results. By pursuing this approach, we will endeavour to solve the same ill-conditioned problem but in the presence of noise.

Contents

1	Introduction	5
1.1	CERN: its history and accelerator complex	5
1.1.1	Historical background	5
1.1.2	CERN complex accelerators	6
1.2	Luminosity	7
1.3	Goals of this thesis work	9
1.4	Thesis Overview	10
2	Mathematical model and computation of the number of collisions	13
2.1	Mathematical model	13
2.1.1	LHC filling schemes	13
2.1.2	Mathematical model	19
2.2	Python implementation	20
3	Maximization of the number of collisions	25
3.1	Formulation of the optimization problem	25
3.2	Unrealistic cases	27
3.2.1	1 st case: Abort Gap, free disposition of the bunches	27
3.2.2	2 nd case: detectors symmetric, SPS batches	28
3.3	Realistic cases	32
3.3.1	3 rd case: Real filling scheme without INDIV	34
3.3.2	4 th case: Real filling scheme with INDIV	36
3.4	Conclusions and following studies	41
4	Mathematical model: luminosity from emittance	43
4.1	The beam emittance and its importance	43
4.1.1	The definition of beam emittance	44
4.1.2	The importance of emittance	47
4.1.3	Emittance measurement	48
4.2	Luminosity model computation	49
4.2.1	The luminosity integral	49
4.2.2	Introducing the parameters	50
4.2.3	Model and hypothesis	53
4.2.4	Cost computation optimization	56
5	Inversion problem: emittance from luminosity	59
5.1	Non-linear system to obtain the emittance	59
5.1.1	The reason why it has been used the luminosity	59
5.1.2	The idea behind this approach	60
5.2	1st inversion problem, without errors	62

5.2.1	1D case : analytical solution	64
5.2.2	2D cases	65
5.2.3	4D case: realistic case	71
5.3	2nd inversion problem, with random errors in the Luminosity output	75
5.3.1	Adding the errors in the approach used before: 1D case	76
5.3.2	2D cases	77
5.3.3	Approaches to overpass this sensitivity	78
5.3.4	4D case: realistic case	83
5.4	Conclusions and following studies	89
Conclusions		92
A		93
A.1	Studies on Long-Range interactions	94
A.1.1	Using the mathematical model...	94
A.1.2	... in a practical point of view	94
A.2	Optimized algorithm Pre-fill without INDIV	95
Bibliography		99

Chapter 1

Introduction

1.1 CERN: its history and accelerator complex

1.1.1 Historical background

The European Organization for Nuclear Research, commonly known as CERN, is a renowned worldwide scientific research organization situated on the border of Switzerland and France. The organization was established in 1954 and currently has a membership of 23 countries. Initially, its primary objective was to understand the composition of atomic nuclei, with a focus on nuclear physics. However, as CERN's energy capacity progressively expanded, its research emphasis transitioned towards the field of particle physics. Scientists at CERN engage in the study of matter's constituents and underlying forces to gain a comprehensive understanding of the fundamental principles governing our Universe. According to the publication by ([CERN]), the Organization adheres to the principle that its research and scientific discoveries shall be openly accessible to the public and will not be influenced by any military obligations.

In addition to its advancements in high-energy physics, CERN has significantly contributed to various academic disciplines. In the latter part of the 1970s, a technique was devised to employ hypertext as a means of arranging and disseminating knowledge within scientific networks. The introduction of hyperlinks facilitated the accessibility of information, leading to the development of the initial web browsers and servers. This significant advancement had a pivotal role in shaping the progression of the World Wide Web, which was subsequently made available to the public in 1990. The discovery of the Higgs boson, which was postulated by Higgs in the 1960s as an explanation for the mechanism by which particles obtain mass, is a notable achievement of the recent years. The Nobel Prize in Physics 2013 was awarded jointly to François Englert and Peter W. Higgs “for the theoretical discovery of a mechanism that contributes to our understanding of the origin of mass of subatomic particles, and which recently was confirmed through the discovery of the predicted fundamental particle, by the ATLAS and CMS experiments at CERN Large Hadron Collider”. Since then, LHC is still the largest circular particle accelerator ever built.

The Organization persistently aims to expand the frontiers of scientific knowledge and technological advancements by employing data analysis, advanced computing systems, and the development of sophisticated engineering technologies. These efforts are directed towards investigating fundamental inquiries pertaining to the nature of the Universe.

1.1.2 CERN complex accelerators

At CERN, the primary instruments used to investigate the forces and components are the accelerators and decelerators, which are designed in two distinct configurations:

- Linear colliders are scientific instruments that employ electric and magnetic fields to drive two particle beams along a linear trajectory, spanning the entire length of the accelerator. The accurate determination of the collision shape is possible due to the singular location where the particles interact. These are single-passage machines, since the particle covers the full accelerator length only once.
- Circular colliders adopt radio-frequency (RF) cavities to impart energy increments, while magnetic fields guide the two particle beams along a circular trajectory. The particles travel along the accelerator ring several billions of times, which increases the likelihood of undergoing rare particle interactions. The interaction points refer to specific positions along the accelerator where the two counter-rotating beams of the collider cross each other, and hence collisions may occur.

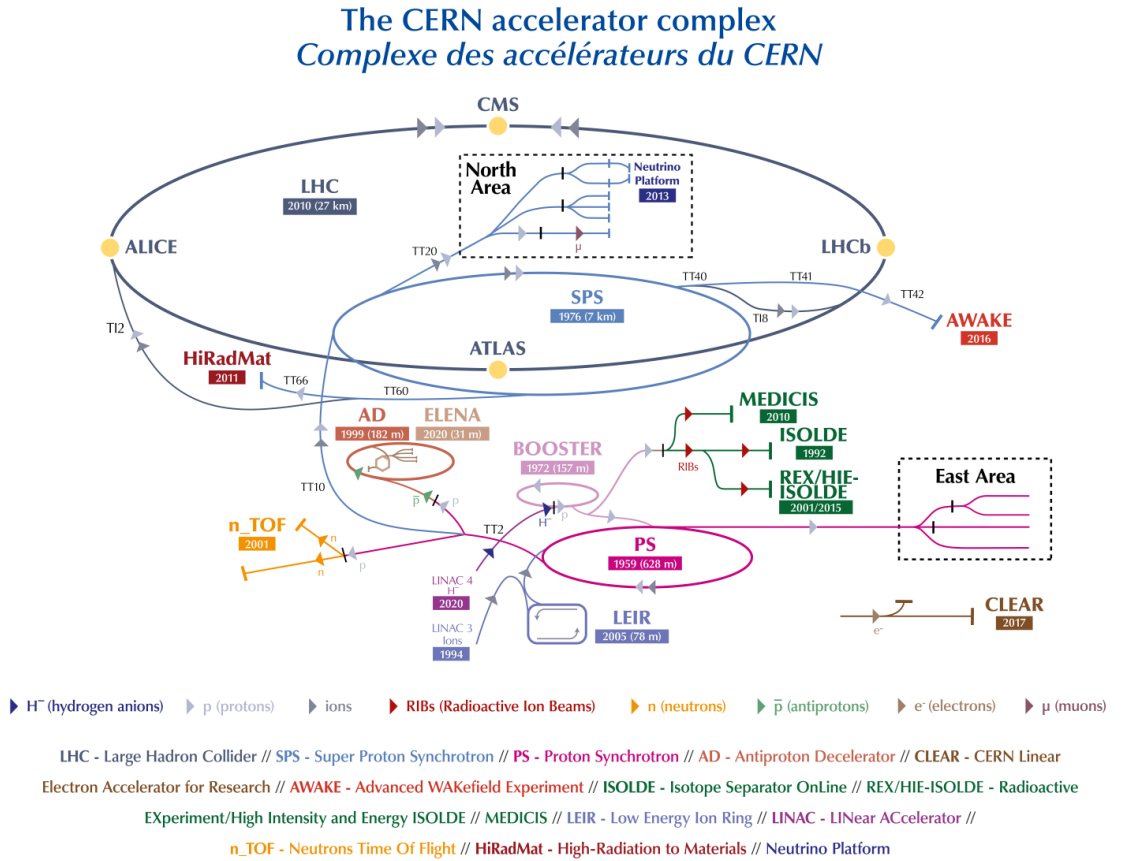


Figure 1.1: CERN complex accelerator in 2022 ([Hal11])

At CERN, a variety of accelerators are used to generate particles with precise energies required for several experimental purposes. These accelerators, which are depicted in Figure 1.1, play a crucial role in facilitating the research conducted at CERN. Particles are subjected to acceleration or deceleration processes to attain different energy levels before their

injection into the subsequent accelerator within this sequence. The sequential use of acceleration phases facilitates the acquisition of the necessary energy to conduct investigations in the field of fundamental physics. As an example, the aforementioned Higgs boson has a mass of $\approx 126 \text{ GeV}/c^2$, that is approximately 130 times larger than the mass of a proton. This boson can be observed in a proton machine only by colliding high energetic protons, with kinetic energy several thousands times their rest mass energy.

Specifically, the supply of protons originates from LINAC4 and thereafter traverses the Protons Synchrotron Booster (PSB), the Proton Synchrotron (PS), and the Super Proton Synchrotron (SPS) before reaching the Large Hadron Collider (LHC) that accelerates them up to 6.8 TeV (that is more than 7000 times their rest mass of $\approx 0.938 \text{ GeV}$).

1.2 Luminosity

To allow for RF acceleration, particle colliders exhibit a specific longitudinal beam structures, wherein particles are organized into several clusters, referred to as bunches, with their quantity varying from a single one to several thousand. In each axis it is possible to define a scalar quantity that represents the arrangement of particles within the bunches along that particular direction. These scalar quantities, are known as emittances. Multiple bunches of particles of the two colliding beams, are accelerated in opposite directions and then collide at a designated interaction point (IP). The detectors, which are positioned at the interaction point, are specifically designed to capture and analyse a vast amount of data resulting from these collisions such as secondary particles, potential instabilities, and other pertinent phenomena.

As illustrated in Figure 1.2, the LHC beams run in two different dedicated pipelines contained within individual vacuum chambers. Even if in each sector, there exists one interior and one external pipeline these overlap near each detector. Indeed, the selection of detector placements is conducted with the intention of maintaining a consistent total circumference for both beams at all times. Each detector is dedicated to the exploration of a distinct branch of physics. These detectors exploit the common physical phenomenon of particle collisions, although with the aim of investigating diverse research domains:

- The ATLAS (A Toroidal LHC ApparatuS) ([AT08]) and CMS (Compact Muon Solenoid) ([CM08]) detectors have been designed specifically to quantify the characteristics of secondary particles generated as an outcome of high-energy collisions. From the collision a substantial amount of data, estimated to be in the range of petabytes, is collected and analysed by a group of experts and scholars. The discovery of the Higgs Boson has been achieved through the utilisation of a specific methodology and the implementation of particular detectors ([Aad12]).
- The LHCb (Large Hadron Collider beauty) experiment ([LH08]), aims to explore rare decays and particle interactions with the purpose of elucidating several enigmas associated with antimatter.
- The ALICE (A Large Ion Collider Experiment) ([AL08]) primarily concentrates on investigations related to heavy ions.

Researchers are currently addressing unresolved inquiries regarding the origins of the cosmos through the examination of particle interactions at elevated energy levels within the LHC.

For further exploration of the LHC experiments, one may refer to the work by [Hal11].

The luminosity, denoted as $\mathcal{L}$, quantified in units of inverse barns per unit of seconds (fb^{-1}Hz), holds significant importance as a key performance indicator in all colliders. The

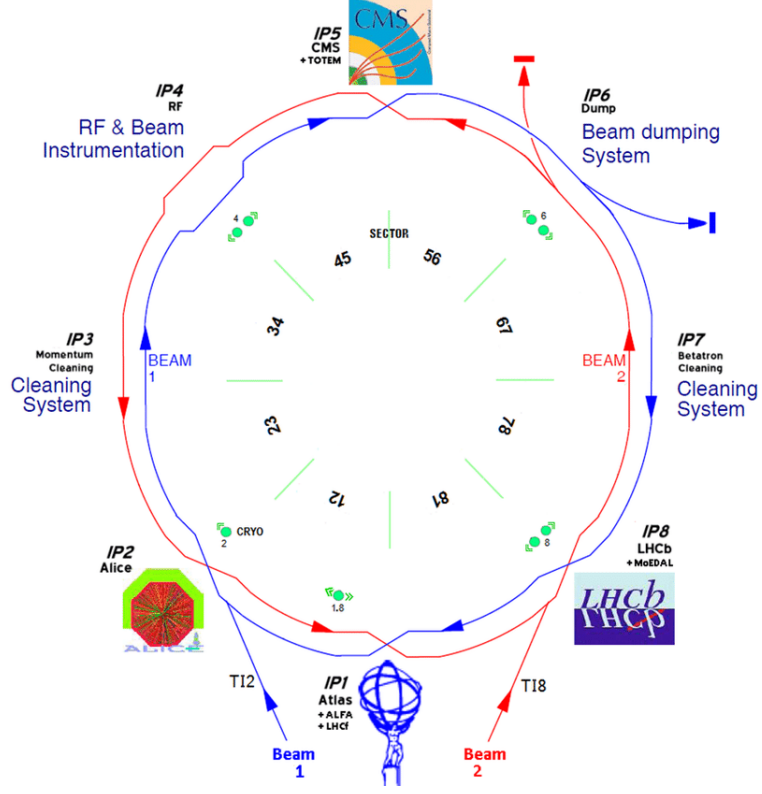


Figure 1.2: Schematic layout of the LHC rings, arcs, and interaction points ([We16])

aforementioned quantity denotes the proportionality between the cross-sectional area of interaction and the rate at which events occur, and it is different in each detector as shown in Figure (1.3)

As a result of enhanced luminosity, the potential for observing a larger number of uncommon physics events is raised, which generate a direct correlation with the collider’s capability for advancing scientific knowledge in the field of physics. The number of bunch collisions for each IP, denoted as $n_{b,IP}$, is strongly linked to the longitudinal arrangement of the bunches in the two beams, referred to as filling schemes, and has a significant impact on the luminosity ($\mathcal{L}$). Indeed, the luminosity for each given IP can be factorized by calculating the cumulative luminosity resulting from all the collisions that occurred at that specific location (notation inspired to the one of the LPC, LHC Programme Coordinator):

$$\mathcal{L}_{IP} = n_{\{bb,IP\}} \mathcal{L}_{bb}. \quad (1.1)$$

where the calculation of $\mathcal{L}_{bb}$ involves computing a quadruple integral over three spatial dimensions and time of the distributions of the two colliding bunches (beam profiles). On the other hand, the determination of $n_{\{bb,IP\}}$ is achieved through discrete convolution of the filling scheme. The filling scheme represents the positions of these bunches within the ring as two boolean vectors, where “True” denotes the presence of a bunch and “False” indicates the absence of a bunch.

To improve the performance of the collider’s machine, it is necessary to optimize the integral of luminosity with respect to time ($\int \mathcal{L}_{IP} dt$) for each experiment. From this, considering that these two factor in the definition (1.1) are both definite positive and independent of each other, it’s possible to optimize them separately.

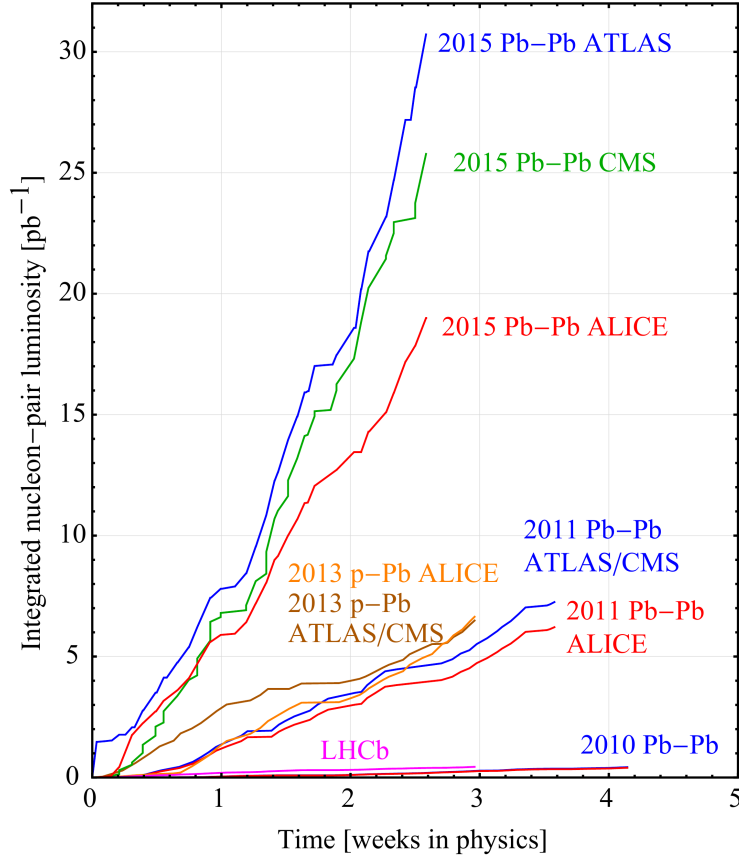


Figure 1.3: Values of luminosity for ions studies, for all the detectors in the LHC, for the first weeks of the RUNs in different years([Wi17])

1.3 Goals of this thesis work

As introduced previously, the luminosity for each given IP can be factorized by calculating the cumulative luminosity resulting from all the collisions that occurred at that specific location, as shown in equation (1.1)

The goal of this thesis is to optimize the luminosity in all its complexity, properly exploiting this factorization.

Firstly, the initial element under analysis pertains to the number of collisions. The objective is to identify a discernible pattern or a rule that governs the arrangement of the bunches. The final purpose is to determine the optimal configuration that maximises the number of collisions throughout all the detectors. The first part of this analysis will introduce a model for determining the number of collisions resulting from the arrangement of bunches. The following part will concentrate on defining the various constraints associated with the placement of the bunches. Once the optimization problem has been comprehensively outlined, the emphasis will shift towards identifying patterns for arranging the bunches. In cases where finding such patterns are not so trivial, strategies will be employed to overcome any encountered challenges.

Next, we will examine the second factor, namely the luminosity bunch-by-bunch, with the goal of using this scalar quantity to derive fundamental parameters known as transverse

emittances (for the x-axis and y-axis). The first part of this analysis will provide a definition of emittance and an explanation of the tools used to measure this parameter in an accelerator. Subsequently, a model will be presented for calculating the luminosity bunch-by-bunch, using a set of distinct parameters known as the configuration system, highlighting the dependence of the emittance within this formula. The second part will be focused on inverting the luminosity model to obtain the desired unknowns. For this reason, after defined the usual configuration system, some of its parameters will be perturbed to generate varying values and model of luminosity bunch-by-bunch. The difference between these luminosity values and the model will be used to establish a system of equations where the emittances are the unknowns, which will then be solved through the Non-linear least squares method. This problem it has been also studied for the case in which the luminosity value is affected by a random error.

1.4 Thesis Overview

The structure of the thesis is as follows:

Chapter 1 is a background chapter

This chapter presents the mathematical model used to calculate the number of collisions. It provides a concise elucidation of the physics principles and notation employed in the arrangement of the bunches, named filling scheme. The emphasis is placed on distinguishing between bunches and RF buckets, differentiating Head-on and Long-range collisions, and highlighting specific accelerator requirements that inherently influence the arrangement of the bunches throughout the accelerator ring. Once defined the mathematical model for calculating the number of bunch colliding in a particular detector, it is presented an optimised implementation in Python, along with a comparison to other possible implementations.

Chapter 2 is a result chapter

This chapter uses the mathematical model and the requirements of the accelerator outlined in Chapter 1 to define an optimization problem. Subsequently, an analysis is conducted on several simplification cases of the optimization problem with the aim of identifying a rule for arranging the bunches in the accelerator ring. The final case under analysis pertains to a general scenario where no discernible patterns or rules are found. Instead, a Monte Carlo simulation is employed, starting from an input filling scheme. This simulation uses the degree of freedom in the arrangement of the bunches, allowing for the exploration of alternative schemes. This chapter concludes with the presentation of interesting results derived from the simultaneous used. These results not only validate certain assertions made in earlier cases examined, but also identify alternative arrangements that outperform the most employed filling scheme in the LHC in 2022. .

Chapter 3 is a background chapter

This chapter provides an exposition of the mathematical model used in the calculation of luminosity bunch-by-bunch. The chapter begins by providing a definition of particle emittance and of beam emittance. Emphasis is placed on the significance of obtaining this parameter, and an examination of the instruments employed thus far to measure it with the greatest possible accuracy. Subsequently, the computation of the luminosity model is delineated, starting with the most general formula. Subsequently examining the distinctive characteristics of the various parameters and the assumptions introduced, the model culminates in the ultimate formula that will be employed in Chapter 4. Furthermore, the relationship between the emittance and the expressed luminosity model is highlighted. Finally,

it is presented an implementation in Python, using Numba, that optimize the computational time of a script already used

Chapter 4 is a result chapter

This chapter employs the mathematical model defined in Chapter 3 to derive the transverse emittances from the luminosity. The reason for selecting luminosity and not other observables of the LHC is elucidated, highlighting the idea and the inspiration underlying the adoption of different configuration systems to capture collisions from multiple perspectives in order to obtain the emittance, placing attention on the initial configuration system. It has been delineating the shifted parameters, the methodology behind, and distinct hypotheses associated with this approach. Subsequently, the analysis progresses from simplified scenarios to more intricate ones, assuming the accuracy and absence of errors in the values obtained from the luminosity model. The concluding section of this study discusses the aforementioned cases, but taking into account random error in luminosity values obtained from the model. The objective in this part is to elucidate the sensitivity of the employed approach and explore alternative strategies for obtaining the different emittances.

Chapter 2

Mathematical model and computation of the number of collisions

2.1 Mathematical model

As reported in Chapter (1.2), there exists a direct correlation between the luminosity of a collider and its potential for making significant physics discoveries. Luminosity is an aggregated information that encompasses several properties of the collider, such as optics and beam dynamics, and produces a scalar value that succinctly represents the effectiveness of the collider.

One of the fundamental dependencies within the mathematical framework underlying the conventional calculation of luminosity is a direct proportion to the number of bunch collisions occurring at the interaction point (IP). For this reason, modern colliders feature several hundred of bunches.

This chapter aims to present the mathematical theory that forms the foundation for the computation of the number of collisions. In addition, the optimized Python code will be analysed, and its results will be compared to the previous computation conducted in CERN ([FillPatt]).

To start with the description of the mathematical model applicable to a broad range of accelerators, it is pertinent to provide a brief overview of the aims and structure of the LHC, hence facilitating the utilization of this model for this particular accelerator.

2.1.1 LHC filling schemes

Definition of bunches

The LHC holds the distinction of being the largest and most powerful circular particle accelerator in the world, boasting a circumference of 27 km. The procedure to accelerate protons or ions to velocities that approach the speed of light entails employing RF cavities and magnets. The acceleration process requires the longitudinal structure of two beams not to be continuous, but rather consists of particles arranged into numerous clusters referred to as bunches, as represented in Figure (2.1). This is because in the event of a continuous beam, the RF cavity induces different acceleration and deceleration of distinct portions of the beam depending on the specific RF voltage represented in the Figure (2.1).

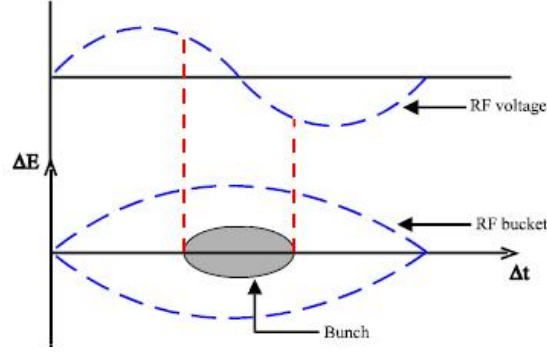


Figure 2.1: The structure of bunches and buckets, considering the Radio Frequency voltage ([Vidal])

The concept of a bunch in accelerator physics is closely interconnected with that of a bucket, which holds significance in understanding the mathematical model:

- Buckets are virtual positions in the LHC ring, corresponding to discrete phasing with respect to the RF cavity, wherein particles receive acceleration or deceleration. In fact, the cavities function at a specific frequency and each period of the RF field defines an RF bucket. The LHC RF cavities function at a frequency of 400.79 MHz, which equates to a **period of 2.5 ns for each LHC RF bucket**. The synchronization of the buckets within the RF cavities plays a crucial role in the acceleration of the particles. Indeed, a necessary requirement for RF cavities is that the ratio, denoted as h_{RF} that is the total number of available buckets, between the revolution frequency of the beam in the ring and the frequency of the cavity must be an integer. This condition ensures that, after completing one full cycle, the cavity returns to phase alignment with the ring's frequency. In the LHC this ratio is 35640, hence LHC has 35640 RF buckets.
- Bunches refer to clusters comprised of various particles within a single beam, sharing a common RF buckets. In each of the three directions, the distribution of the particle in the bunch is characterized by a particle density, the transverse beam profiles are modelled as Gaussian, the longitudinal beam profile is depicted in Figure (2.2). These bunches are established during the initial stages of the LHC chain as particles undergo energy modifications, after which they are subsequently gathered using various methods. Bunches play a vital part in the efficiency of an accelerator due to their critical beam properties, such as the beam profile, and the consequential number of collisions resulting from their placement within the accelerator.

From a technical point of view, it is possible to establish a one-to-one connection between the bunches and buckets because each period RF bucket that may hold a particle bunch, this would imply that bunches might also be evenly spaced at intervals of 2.5 nanoseconds. Nevertheless, in this particular arrangement, the detectors would encounter limitations in their capacity to handle a substantial amount of data, as well as their inability to differentiate between recently generated secondary particles and those that have been present for a longer duration. This was one of the reason why it was decided to **arrange protons** in clusters at intervals of 7.5 meters, one bunch **every 10 RF buckets (25 nanoseconds)**.

For further insight into the process of aggregating various bunches in the accelerator chain of the LHC, relevant literature sources such as ([Da18]) and ([Ga01]) provide valuable information.

Based on this concise elucidation, it can be asserted that each bunch can be injected in every 10th bucket. This injection bucket is denoted as bunch slot.

In the following, “slot” and “bucket” will be considered synonyms.

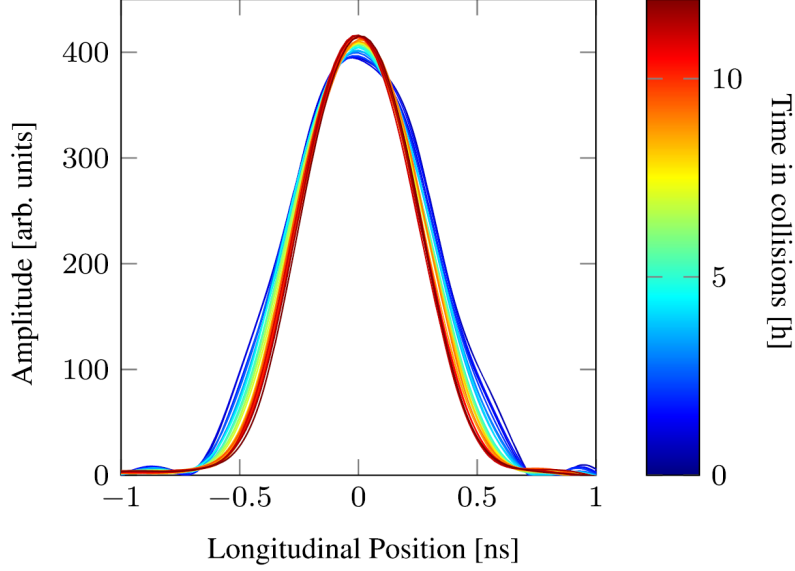


Figure 2.2: Longitudinal beam profiles over 12 hours in collisions ([Hos18])

Filling scheme

The filling scheme, as seen in Figure (2.3), refers to the specific sequence in which the bunches are injected and circulated within the ring to attain a specific level of machine performance. From this definition, it is easy to observe that the total length of the filling scheme in the LHC is **3564** (35640/10) slots. Several filling schemes have been implemented in the LHC to accommodate the diverse requirements and outcomes of the accelerator. In particular, several elements must be taken into account during the development of a filling scheme:

- **Machine structure:** The LHC represents the final stage of the accelerator complex at the CERN. Within each ring of the accelerator complex, the particle beam undergoes acceleration before being injected into a bigger ring. The interconnection between two consecutive rings, as shown in Figure (1.1), is determined by the application of radio frequencies (RF), which meticulously control the arrangement of particles in the LHC filling scheme.

These cavities play a crucial role in ensuring that the spacing between each bunch is precisely maintained at 25 nanoseconds. This specific bunch spacing is achieved by the Proton Synchrotron (PS), which gathers the various bunches into a batch known as the PS batch. Typically, the PS batch consists of a multiple of 12 bunches, each separated by 25 ns. The **PS batches** are injected into the SPS (Super Proton Synchrotron) to generate the **SPS batch** (also referred to as ‘train’), which consists of multiple PS batches spaced by minimum 200 ns (**7 slots**) as determined by the RF system. This batch is then accelerated from an initial energy of 26 GeV to a final energy of 450 GeV. Subsequently, the aforementioned trains are introduced into the LHC, with a minimum spacing of 800 nanoseconds between them (equivalent to **31 slots**), with

the purpose of achieving acceleration up to 6.8 teraelectronvolts (TeV) per beam. For further understanding of the functioning of the RF system in the LHC, readers are encouraged to refer to ([Bo99]). Additionally, for more comprehensive information regarding the various spacing and gaps inside the LHC ring, ([Ve17]) provides detailed insights.

- **Machine Diagnostic:** The filling system also takes into account certain needs pertaining to diagnostic procedures. Specifically, the presence of non-colliding bunches is necessary for the purpose of machine and detector diagnostics, despite its impact on the number of collision and subsequent luminosity performance. One of the most crucial measures is related to the “betatron tune” (q) which is closely associated with the oscillation of the beam within the transverse plane. This tune is measured specifically using the initial **twelve bunches**, which are intentionally **prevented from colliding within the ATLAS/CMS detectors**. For further elaboration on the betatron tune, interested readers may refer to the work ([Jo18]) for additional information.
- **Machine protection:** it is a crucial element for the preservation of the LHC, as it is a very intricate accelerator that necessitates the implementation of numerous safety protocols. Among these protocols, the Abort gap stands out as an important one.

The procedure of dumping the beams of the LHC is a crucial and intricate operation. It must be executed with great caution to prevent magnet quenching, damages, and minimize losses during the rise time of the LHC extraction kicker. To accomplish this, a $\approx 3 \mu s$ (**121 slots**) gap (i.e. without bunches) known as “**abort gap**”(AG) is positioned in the filling scheme. This gap ensures that a portion of the LHC ring remains free of particles during the transient related to the dumping process. The initial RF bucket located beyond the abort gap is designated as bucket 1. To ensure the absence of bunches within the abort gap, the final slot eligible for bunch train injection is determined by adding the length of the abort gap to that of the longest previously injected bunch train. The sum of the abort gap and the longest bunch train of the filling scheme defines the so-called “**Abort Gap Keeper**” (AGK).

For a more comprehensive understanding of the Abort Gap and its associated mechanism, interested readers are encouraged to refer to the work ([We17]).

- **Disposition of the detectors:** As shown in the Figure (1.2) the LHC is organized into eight distinct octants, with each octant accommodating a total of **445.5 slots**. The positions of the detectors in the LHC can be denoted in terms of the longitudinal variable, s , relative to the circumference of the LHC ($C \approx 27 \text{ km}$):

$$\text{ATLAS (IP1): at } s = 0, \quad (2.1)$$

$$\text{ALICE (IP2): at } s = \frac{C}{8}, \quad (2.2)$$

$$\text{CMS (IP5): at } s = \frac{C}{2}, \quad (2.3)$$

$$\text{LHCb (IP8): at } s = \frac{7C}{8} - \frac{15C}{h_{RF}}. \quad (2.4)$$

It is evident that, except the **LHCb** which is **displaced** by a distance of 11.22 m (equivalent to **1.5 times** the distance between consecutive slots) from the centre of the eighth octant, the majority of the interaction points are located in the centre of the octants. Therefore, the LHC does not exhibit an octagonal symmetry. Indeed, in the

event that we introduce four equidistantly distributed bunches in the same positions for the two beams, this will result in a collision occurring at IP1, IP5, and IP2, but not in IP8.

- Experiment requirements: The filling scheme considers various detector needs, such as collision count and bunch spacing. One of the most intriguing obligations to fulfil within the context of the LHC involves the **INDIVs** (individual bunches). These trains, consisting of a solitary bunch, **are employed for collision** purposes with another INDIV **within a designated detector**. This specific collision event is used for detectors' setup and calibration.

All these aspects affect the structure of the filling scheme, as represented in the Figure (2.3).

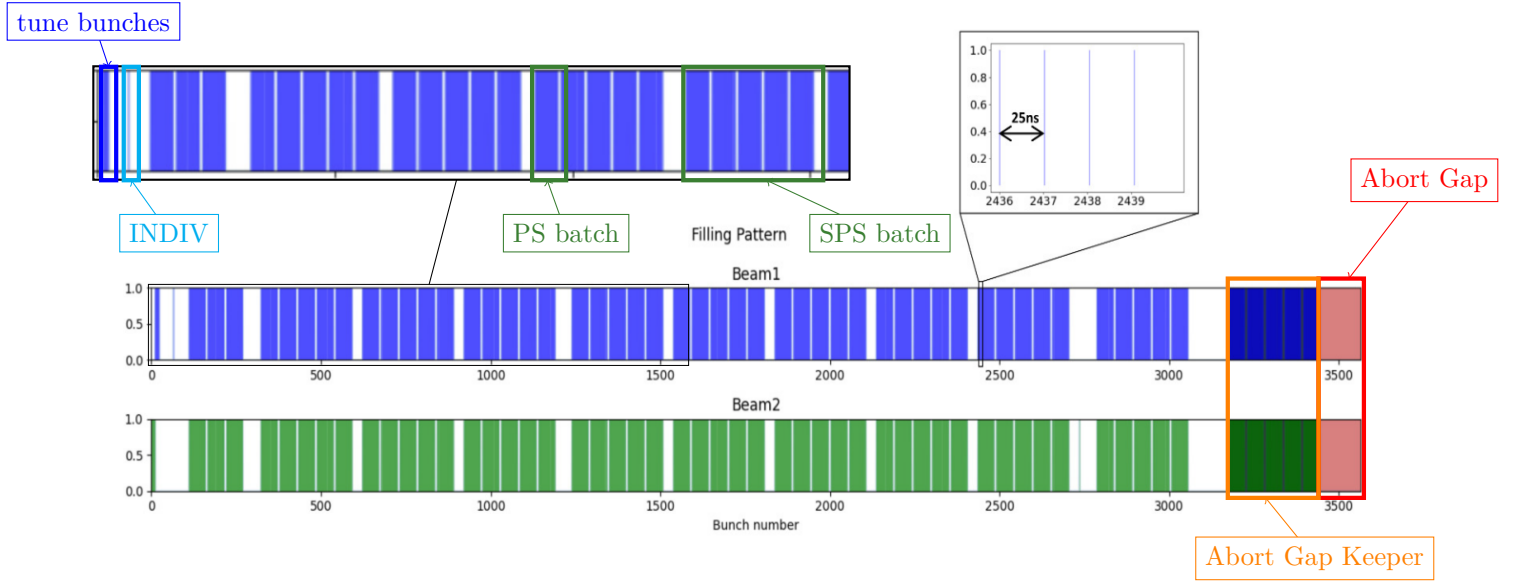


Figure 2.3: Example of filling scheme used at CERN, B1 in blue and B2 in green. In the zoom, it is represented the distance between bunches that is lost from the filling scheme representation. In different colours are represented, how the various aspects presented affect the filling scheme structure.

Due to the inherent complexity of the accelerator ring, acquiring information regarding the collisions is a non-trivial task.

Collisions in LHC

The computation of a collision schedule involves determining the specific bunches and their collision locations. This computation relies on the positions of the detectors and the two filling schemes employed to represent the beams. The types of collision, as represented in Figure (2.4), can be split into:

- Head-on collisions, commonly referred to as “central collisions”, occur at the **centre of the detector**. In this type of collisions, the centre of mass energy is maximized as a result of the physical collision of the beams. This form of collision has the potential to generate novel particles, while also enabling the observation of rarer particles.
- Long-range collisions occur within a common vacuum chamber in **proximity to the interaction point**. Despite being physically separated (no nuclear interactions), the

two beams experience **electromagnetic forces** generated by the other one. Hence, these types of collisions are commonly referred to as “parasitic”. The number of long-range collisions at each interaction point is contingent upon the threshold longitudinal distance at which one may disregard electromagnetic collisions between the two beams. The threshold in question pertains to the accelerator’s facilities. Indeed, in the LHC at a distance of around 60 meters (equivalent to 8 slots), there exists a dipole. Beyond this dipole, the separation between the beams rises rapidly. Hence, this quantity is not predetermined, but rather contingent upon the user’s purpose and the configuration of the accelerator. If the user is interested in studying dipolar (linear) effects or higher order effects, he may need to take into account the distance between the magnet and the interaction point in order to determine the appropriate number of long-range collisions to investigate.

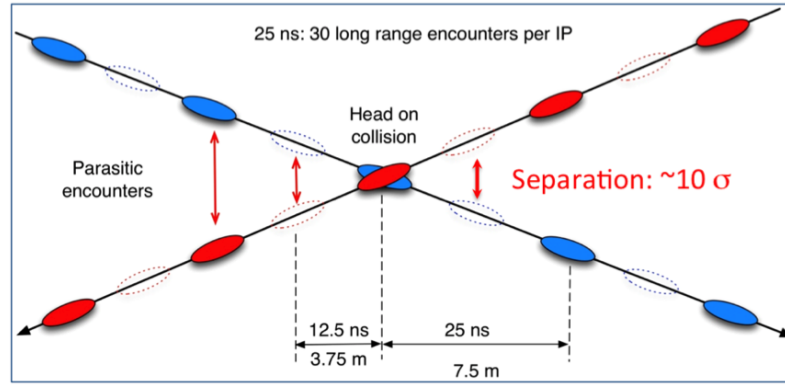


Figure 2.4: Representation of Head-On collisions and Parasitic ones, ([La16])

For additional information, please refer to the publication by W. Herr([He01]).

It would be of great interest determining an optimal approach for the computation of a detailed schedule for both types of collisions in a general accelerator, with particular emphasis on the LHC due to its significant number of bunches (named or collision schedule or BB encounter schedule). This schedule should encompass the following elements:

- the data regarding the **number of HO and LR collisions** occurring at an interaction location,
- the number of respective **collision partners**, starting from bucket 1, in either the HO and LR collision at a given interaction point for each bunch (named BB partner),
- the **distance** from the interaction point, if the collision is **long-range**.

Furthermore, due to the presence of accelerator constraints, the bunches are not evenly spaced along the LHC-ring, but instead present distinct intervals or gaps. Consequently, the particles experience different beam-beam forces, leading to an intricate collision schedule. Establishing a clear and concise model is of the highest priority in order to effectively compute all necessary details pertaining to the collision schedule. Moreover, it would be highly advantageous if this model could be used to both Head-On and Long-Range collisions.

Before looking into the model, it is important to address a minor detail. The LHC is a particle accelerator consisting of two rings, one interior and one external, each containing an isolated vacuum chamber. Within these chambers, two beams of particles circulate at a synchronized revolution frequency. It is important to note that while the vacuum chambers

may have shared segments and the radio frequency (RF) may need the total circumference of the two beam routes to be equal, it is possible for the length of the two beam paths to vary from a specific reference point to a more generic one. However, in our analysis, which is grounded in the findings of ([Jo99]), we can assert that the paths covered by two beams in a single revolution are identical.

2.1.2 Mathematical model

The aim of this study is to develop a computational model that accurately represents the BB encounter schedule.

We aim to obtain a data frame as output of this calculation, which will encompass comprehensive details regarding the number of Head-On collisions taking place at a specific ring slot (an interaction point). Additionally, the data frame will include information about the partner involved in each collision (HO or LR), as well as the number of Long-Range interactions occurring with different bunches and the corresponding locations relative to an interaction point, here are reported some columns and rows of this data frame for LHCb.

index, # of LR in LHCB, HO partner in LHCB, BB partners in LHCB,	Positions in LHCB
135, 9.0, nan, [2817.0, 2818.0, 2819.0, 2820.0, 2821.0, 2822.0, 2823.0, 2824.0, 2825.0],	[12.0, 13.0, 14.0, 15.0, 16.0, 17.0, 18.0, 19.0, 20.0]
136, 10.0, nan, [2817.0, 2818.0, 2819.0, 2820.0, 2821.0, 2822.0, 2823.0, 2824.0, 2825.0, 2826.0],	[11.0, 12.0, 13.0, 14.0, 15.0, 16.0, 17.0, 18.0, 19.0, 20.0]
137, 11.0, nan, [2817.0, 2818.0, 2819.0, 2820.0, 2821.0, 2822.0, 2823.0, 2824.0, 2825.0, 2826.0, 2827.0],	[10.0, 11.0, 12.0, 13.0, 14.0, 15.0, 16.0, 17.0, 18.0, 19.0, 20.0]
138, 12.0, nan, [2817.0, 2818.0, 2819.0, 2820.0, 2821.0, 2822.0, 2823.0, 2824.0, 2825.0, 2826.0, 2827.0, 2828.0],	[9.0, 10.0, 11.0, 12.0, 13.0, 14.0, 15.0, 16.0, 17.0, 18.0, 19.0, 20.0]
1029, 11.0, 135.0, [136.0, 137.0, 138.0, 139.0, 140.0, 141.0, 142.0, 143.0, 144.0, 145.0, 146.0],	[1.0, 2.0, 3.0, 4.0, 5.0, 6.0, 7.0, 8.0, 9.0, 10.0, 11.0]
1030, 11.0, 136.0, [135.0, 137.0, 138.0, 139.0, 140.0, 141.0, 142.0, 143.0, 144.0, 145.0, 146.0],	[-1.0, 1.0, 2.0, 3.0, 4.0, 5.0, 6.0, 7.0, 8.0, 9.0, 10.0]
1031, 11.0, 137.0, [135.0, 136.0, 138.0, 139.0, 140.0, 141.0, 142.0, 143.0, 144.0, 145.0, 146.0],	[-2.0, -1.0, 1.0, 2.0, 3.0, 4.0, 5.0, 6.0, 7.0, 8.0, 9.0]
1032, 11.0, 138.0, [135.0, 136.0, 137.0, 139.0, 140.0, 141.0, 142.0, 143.0, 144.0, 145.0, 146.0],	[-3.0, -2.0, -1.0, 1.0, 2.0, 3.0, 4.0, 5.0, 6.0, 7.0, 8.0]

From the physics behind the phenomenon, it has been decided to describe this computation using a discrete convolution in the longitudinal space. This approach is taken in order to precisely extract the required information and emphasize the similarities between the two categories of collisions. In the following, we will present the hypothesis considered and the final model.

It is postulated that the two beams will engage in collision within the designated slot, wherein the objective is to calculate the number of collisions. Consequently, it can be inferred that the two beams will occupy the same vacuum chamber. Due to this assumption, information pertaining to planes other than the longitudinal plane is considered unnecessary for this computation.

Therefore, the two filling schemes can be represented as boolean vectors, where each slot is assigned a value of 1 if it is filled and 0 if it is empty. This modelling approach is suitable since the filling pattern is not concerned with individual particles, but rather with bunches of particles. Therefore, it is sufficient to consider the positions of these bunches rather than the detailed beam profiles.

Moreover, considering the fact that the two beams are moving in opposite directions, we can determine the number of collisions occurring in a specific ring slot by summing the overlapping bunches generated through the shifting of the two filling schemes:

$$n_{coll}[id_{slot}] = \sum_{i=0}^{\frac{h_{RF}}{10}=3564} B_1[i - \text{mod}(id_{slot}, 3564)] * B_2[i + \text{mod}(id_{slot}, 3564)] = \quad (2.5)$$

$$\sum_{i=0}^{\frac{h_{RF}}{10}=3564} B_1[i - \text{mod}(2 * id_{slot}, 3564)] * B_2[i]$$

where id_{slot} denotes the detector's position in units of bunch slots, whereas the term “mod” represents the module function, as the entire ring exhibits periodicity after completing one full turn.

We are particularly interested in the number of collisions occurring at the four interaction points of the LHC. The positions of the detectors, referred to as the IP pattern, are obtained by dividing the equations (2.1)-(2.4) by the bunch spacing:

$$\left[\underbrace{0}_{IP1}, \underbrace{445.5}_{IP2 \text{ at } \frac{3564}{8}}, \underbrace{1782}_{IP5}, \underbrace{3117}_{IP8 \text{ at } \frac{7*3564}{8}-1.5} \right] \xrightarrow{\text{mod}(2id_{slot}, 3564)} [0, 891, 0, 2670] \quad (2.6)$$

The number of head-on collisions occurring at the LHC can be efficiently and easily determined by using the provided positions. It is important to note that this calculation holds true for both Head-On and Long-Range collisions.

A theorem is presented in Appendix A.1 as an application of the mathematical model to determine filling schemes that in specific slots do not have Long-Range interactions.

2.2 Python implementation

One of the objectives of this mathematical model was to translate it in a Python function, with the aim of speeding-up the calculation process of the beam-beam encounter schedule. Before the approach presented in this work, the software employed at CERN for the beam-beam encounter schedule was suboptimal ([FillPatt]).

Instead of performing a discrete convolution on the boolean vectors, one may easily determine the number of collisions by exploiting the information in the frequency domain through the utilization of a Fast Fourier Transform:

```
import numpy as np
from scipy.fft import fft, ifft
N = 3564
B1 = np.random.randint(0,2,N)
B2 = np.random.randint(0,2,N)
id_slot = np.random.randint(0,N,1)

# Computation in frequency domain
fft_B1 = fft(B1)
fft_B2 = fft(np.flip(B2))
n_coll_freq = np.flip(ifft(fft_B1*fft_B2))
# Computation in time domain
n_coll_time = (np.roll(B1,id_slot)*B2).sum()

assert int(np.abs(n_coll_freq[id_slot]))==n_coll_time
```

While both computations yield equal results in a technical sense, one of them is more computational efficient. The computing cost of the two approaches are as follows:

1. Using the convolution: $N_{ids} * N$

2. Using the FFT: $N * \log(N)$

where N_{ids} represents the quantity of interaction points at which the user want to calculate the number of collisions and N the length of the filling scheme. Subsequently, the relative magnitude of the ratio $\frac{N_{ids}}{\log(N)}$ determines the comparative efficiency of the two methods. It is important to note that $N_{ids} \ll N$, limiting the performance of the FFT method.

A further drawback associated with employing the FFT for determining the number of collisions, is the computation of the partner bunches and the precise location of the LR collision. Their computation can be easily performed in the time domain. Prior to computing the final sum of the convolution in a specific slot, it is possible to determine the information regarding the partner bunches involved in collisions within that slot. This can be achieved by selecting an element from the vector resulting from the product, which has a value of 1. By subtracting (summing) the position of the chosen slot from the index of that vector element, the output value corresponds to the number of the partner bunch for B1(B2) in that specific slot. In the context of frequency domain analysis, it remains uncertain whether the conversion of this calculation is feasible, properly because the information represented in the frequency domain exhibits a high level of complexity. Consequently, it has been determined that adopting the time domain analysis is more advantageous, primarily due to its optimized computing efficiency and also for its simplicity of interpretation. This limitation arises due to the analysis of data in the frequency domain rather than in the time domain.

As a result of these factors, we choose to compile a Python script for the purpose of calculating the BB encounter schedule using the first technique. This implementation resulted in a significant enhancement in computational efficiency, that is represented in the histograms in Figure (2.5).

most used filling scheme in 2023	most used filling in 2022
1178:1178b_1165_696_693_9inj_2INDIV	1551:1551b_1538_1406_1472_12inj_3INDIVs
1818:1818b_1805_1057_1182_14inj_2INDIV	1935:1935b_1922_1602_1672_14inj_3INDIVs
1886:1886b_1873_1217_1173_12inj_2INDIV	2173:2173b_2160_1804_1737_11inj_1INDIV
1903:1903b_1890_1099_1160_12inj_3INDIV	2390:2390b_2378_1967_2106_18inj_2INDIV
2358:2358b_2345_1692_1628_14inj_2INDIV	2413:2413b_2400_1836_1845_12inj_1INDIV
2374:2374b_2361_1730_1773_13inj_2INDIV	2461:2461b_2448_1737_1733_16inj_1INDIV
2464::2464b_2452_1842_1821_12inj	2462:2462b_2450_1737_1735_17inj_2INDIV

Table 2.1: Table representing the notation for most used filling scheme in the LHC used as x-axis in the Figure (2.5), the notation represent “NumberOfBunches_ATLAS/CMS_ALICE_LHCb_NumberOfjections_NumberOfINDIV”. In particular, the filling schemes used in the 2023 are all hybrid filling schemes.

The histograms presented in Figure (2.5) clearly illustrate the disparities between the two algorithms. Specifically, the figure depicting the cumulative time highlights a significant enhancement of approximately 50 times when comparing the newer method to the older one. Conversely, the plot illustrating the total time clearly demonstrates a two-order-of-magnitude improvement between the two algorithms. Another intriguing aspect of the histogram is its ordering of the various filling schemes based on the total number of bunches within the ring. It is observed that both the computation time of the older algorithm increases in accordance with the chosen order. Conversely, the new algorithm exhibits different fluctuations that are not easily interpretable, as well as anomalous decreases, such as the one observed on the right side of the picture subsequent to the first fillig scheme. It can be readily asserted that these fluctuations may be contingent upon several circumstances, such as the number of collisions on a particular detector or the arrangement of the particle

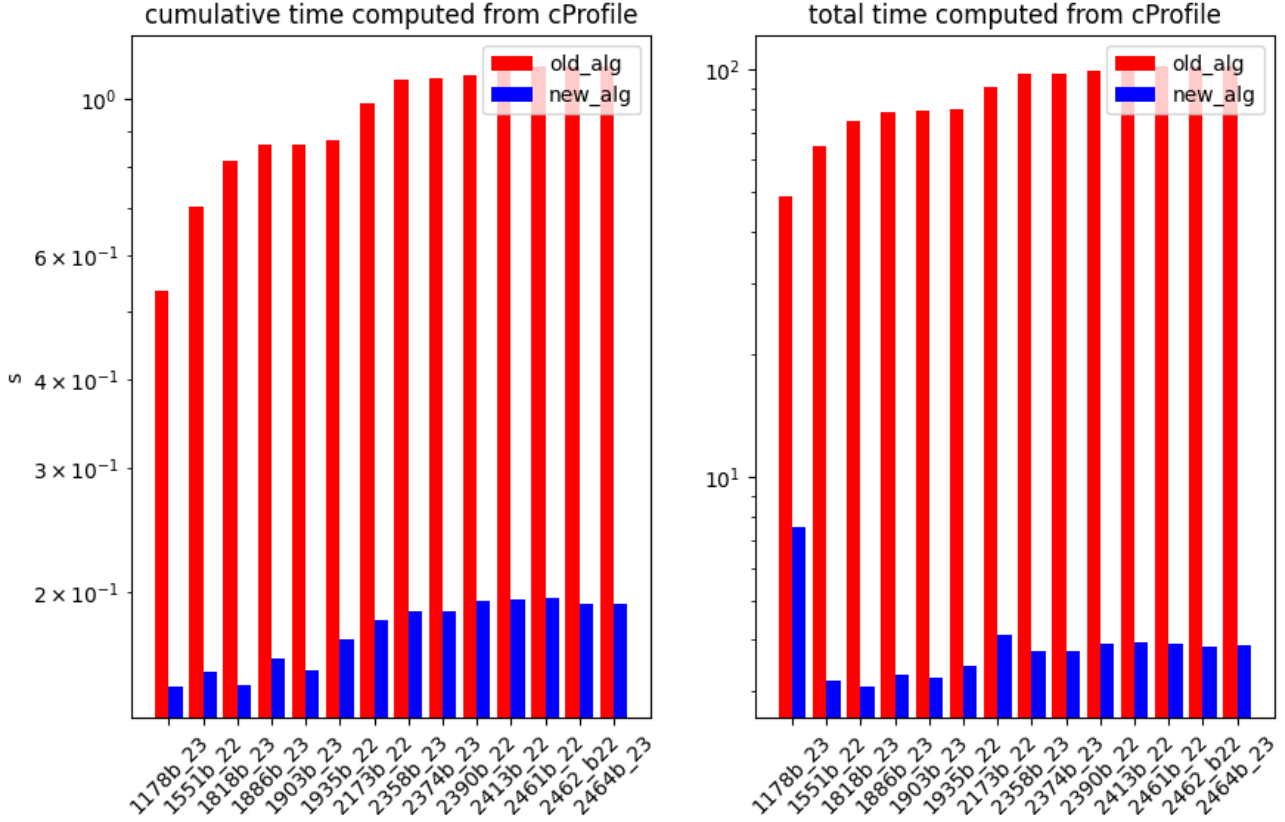


Figure 2.5: Histogram of the performances of the two algorithms tested on the most used filling schemes on 2022 and 2023 in the LHC, shown in the Table (2.1), the order shown in the x-axis is increasing with respect to the number of bunches of the filling scheme. Both the histogram have been plotted from a time average of 10 repeated algorithms on the same filling scheme obtained from cProfile in python. On left is represented the cumulative time, on right is represented the total time.

bunches. However, it is important to note that there exists a counterexample that effectively refutes this assertion. Indeed, it can be noticed that on the left side of the picture there is a significant reduction in computational time between the cases 2461_22 and 2462_22. Although the two filling schemes are nearly identical, with the only difference being that the second one has one INDIV bunch more for each beam, but the arrangement of the bunches remains consistent in the same slots. This behaviour completely disprove that this peculiar behaviour it is governed by the bunch disposition in the filling scheme. This phenomenon could be elucidated, taking into account the utilization of additional libraries such as numpy in the algorithm, by proposing that there is a form of memory inside the framework of algorithmic computation. This indicates that, after repeated testing, the Python compiler preserves certain steps of the method in memory. Consequently, in subsequent tests, the compiler bypasses certain steps in order to enhance efficiency. In order to mitigate this behaviour, we have chosen to recompute the identical computation, while recording the distinct values and restarting the Python kernel either for both the algorithms. The ensuing outcome is presented in the Figure (2.6)

In order to know more details about this implementation, it is possible to see ([CollSched]).

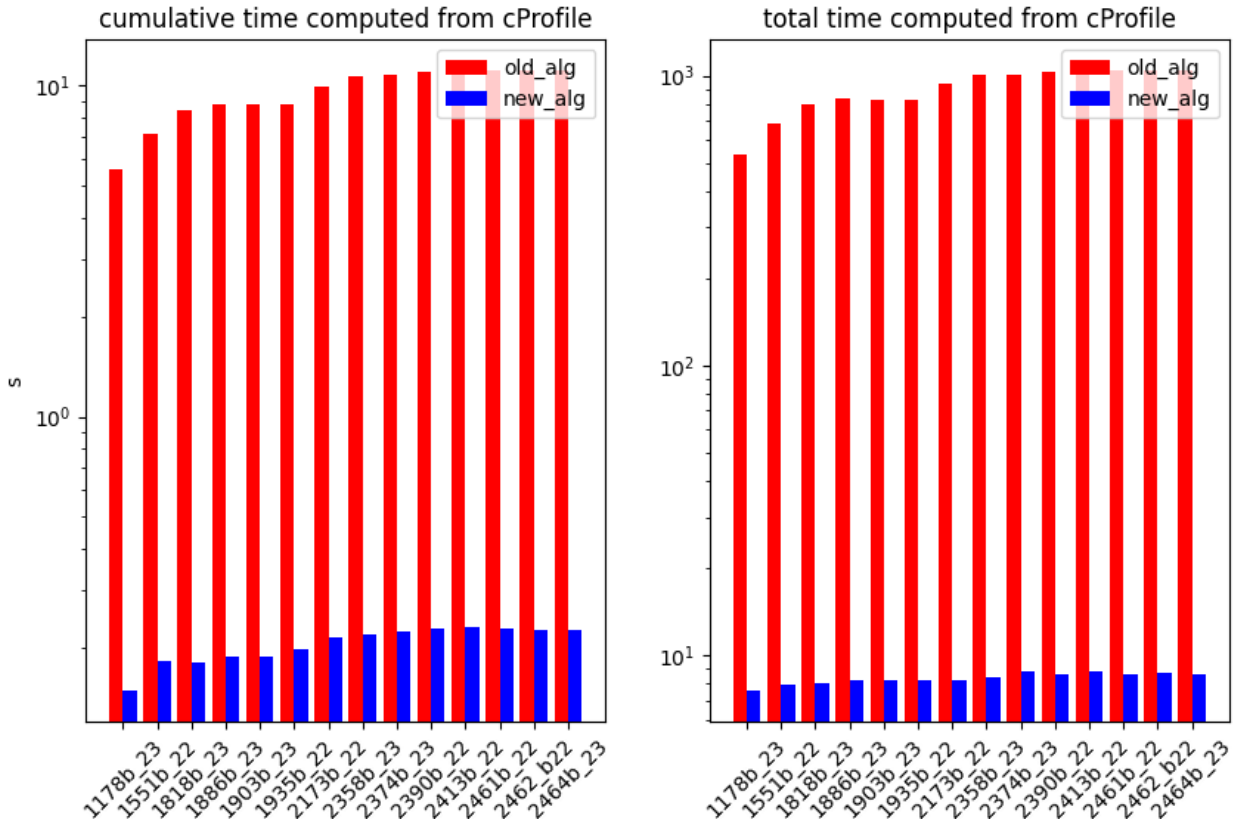


Figure 2.6: Histogram of the performances of the two algorithms tested on the most used filling schemes on 2022 and 2023 in the LHC, shown in the Table (2.1), the order shown in the x-axis is increasing with respect to the number of bunches of the filling scheme. Both the histogram have been plotted from a time average of 10 repeated algorithms on the same filling scheme obtained from cProfile in python, where passing from one filling scheme to another the python kernel has been restarted. On left is represented the cumulative time, on right is represented the total time.

Chapter 3

Maximization of the number of collisions

In the preceding chapter, our objective was to determine, for a given filling scheme, the most effective method for calculating the BB encounter schedule.

In this chapter, the approach is reversed, aiming to study the problem synthetically. The objective is to identify a pattern or rule governing the arrangement of bunches, with the ultimate goal of determining the best filling scheme. The reader may inquire as to which objective it is feasible to establish a superior filling scheme. The subsequent explanation will elucidate this matter.

In order to enhance the operational efficiency and performance of the accelerator, it is imperative to maximize the integrated luminosity at the interaction points over time. In this purpose there are a lot of factors that play a crucial rule, one of them is to arrange the bunches in a manner that maximizes the number of collisions at each interaction point. This chapter will primarily address the task of identifying a governing principle for determining a filling scheme that may optimize the number of collisions in each detector within the LHC. However, it is pertinent to question the feasibility of achieving such an objective.

Firstly, this analysis will present certain features and limitations associated with the LHC, which are also evident in the arrangement of particle bunches. After formulating the comprehensive optimization problem, a systematic analysis will be conducted, starting with simpler examples with different assumptions and progressing, by incremental steps, towards the fully-fledged problem.

3.1 Formulation of the optimization problem

Given the mathematical model offered in the preceding chapter (2.1.2) and considering the limitations and requirements imposed by the filling scheme outlined in the previous chapter (2.1.1), we may now formulate our optimization problem in mathematical terms:

$$\begin{aligned}
 & \max_{B1, B2 \in \mathbf{R}^{3564}} \left. \begin{array}{l} (B1 * B2)[0] \\ (B1 * B2)[891] \\ (B1 * B2)[2670] \end{array} \right\} \text{Real detectors' disposition} \\
 & \text{s.t. } \underbrace{\exists! \text{ PS}(x, 12) \in B1, B2, \text{ where } x < 36}_{\text{Machine diagnostic: bunches for tune measurement}} \\
 & \underbrace{|x_2 - x_1| \geq \text{bun}_q, \text{ for } \text{PS}(x_1, \text{bun}_q)_{B1}, \text{PS}(x_2, \text{bun}_q)_{B2}}_{\text{Machine diagnostic: bunches for tune measurement}} \\
 & \underbrace{\delta(x)_{B1}, \delta(x)_{B2} \in B1, B2, \text{ has specific collisions}}_{\text{Experiment requirements: INDIV}} \\
 & \underbrace{\text{SPS}_{\mathbf{n}}(x_1, \text{bun}) = \sum_{i=1}^{\mathbf{n}} \text{PS}(x_i, \text{bun}), \text{ where } x_{i+1} \geq x_i + \text{bun} + 8}_{\text{Machine structure: RF cavity}} \\
 & \underbrace{\sum_{i=1}^{n_{\text{train}}} \text{SPS}_{\mathbf{n}_i}(x_i, \text{bun}) \in B1, B2, \text{ where } x_{i+1} \geq x_i + \mathbf{n}_i \text{bun} + 32}_{\text{Machine structure: RF cavity}} \\
 & \underbrace{\delta(x), \text{SPS}_{\mathbf{n}}(x, \text{bun}), \text{ where } x \leq \text{AG} = (3564 - 121)}_{\text{Machine protection: Abort Gap and Abort Gap Keeper}}
 \end{aligned} \tag{3.1}$$

where B1 and B2 denote the boolean vectors representing the filling scheme of two beams. The discrete δ -function, denoted as $\delta(x)$, is used to represent the **INDIV**. The discrete door function, $\text{PS}(x, \text{bun})$, is employed to represent the **PS batch**. The train (**SPS batch**) is represented by $\text{SPS}_{\mathbf{n}}(x, \text{bun})$, composed by $\mathbf{n}$ PS batches, where n_{train} refers to the number of injected trains utilized in the filling scheme and bun_q represent the number of bunches for tune measurement. This number is the same for both beams. The indices at which we aim to maximize the convolution are obtained from the module function of the detectors' positions from (2.6). Based on the mathematical representation, it is evident that the maximization of collisions in four distinct detectors can be achieved by using only three indices. This conclusion is derived from the module function incorporated in the mathematical model. Indeed, it is worth noting that due to symmetry considerations, the ATLAS and CMS detectors will consistently exhibit an equal number of collisions.

The mathematical representation of this problem reveals its non-trivial nature in terms of optimization. Firstly, the objective is to maximize the convolution only at specific points that deviate from the 8-fold symmetry mentioned in the introduction. Secondly, numerous constraints must be considered to ensure the feasibility of a filling scheme for the accelerator.

The ultimate objective of our study is to develop an algorithm that, given the length of the empty gap between consecutive PS batches, when provided with specific input as:

- number of bunches per PS batch
- number of INDIV per beam
- maximum length for an SPS batch

- number of bunches per tune measurement
- number of slots to inject the bunches for tune measurement

can determine the most optimal filling scheme satisfying the constraints of the optimization problem.

In each analysed example, the minimum number of vacant slots between two successive PS batches and the number of bunches for each PS batch will be fixed as a constant for all the resolution methods. So the final goal of the approach is to determine the empty space between two SPS batches and their length, properly because the information of the PS batch will be fixed from the user.

To facilitate a comprehensive understanding of the problem at hand, it has been employed a systematic approach wherein the complexity will be gradually escalated in a stepwise manner.

3.2 Unrealistic cases

The objective of our study is to identify the optimal filling scheme, specifically the filling scheme that maximizes the number of collisions across all detectors simultaneously.

From a mathematical perspective this problem is well-defined, but in some configuration it is impossible to solve, as it will be elucidated later. It is important to note that this analysis will be limited to certain simplified scenarios.

3.2.1 1st case: Abort Gap, free disposition of the bunches

In this particular scenario, it is assumed that there are no constraints on the number of bunches that can be injected into the ring. This allows for the ring to be filled without following any established structure, such as SPS batches, PS batches, INDIV, or bun_q. By eliminating all potential structures, the limitations associated with RF cavities for all proton injections in the proton chain are removed. Consequently, there is no longer a need to maintain a vacant interval between the bunches, except the 25 ns gap, which is not visually depicted in the boolean vectors. The sole constraint taken into account in this scenario is the one related to machine protection, but with a slight modification. Defining an abort gap killer (AGK) becomes problematic when the concept of SPS batch is eliminated. Consequently, in this scenario, the focus is solely on the limitation of the abort gap, without considering the AGK.

$$\begin{aligned}
 & \max_{B1, B2 \in \mathbf{R}^{3564}} \left. \begin{array}{l} (B1 * B2)[0] \\ (B1 * B2)[891] \\ (B1 * B2)[2670] \end{array} \right\} \text{Real detectors' disposition} \\
 & \text{s.t.} \quad \underbrace{\forall x \leq AG = (3564 - 121)}_{\text{Machine protection: all the bunches injected before the Abort Gap}}
 \end{aligned}$$

In this particular arrangement, taking into account the actual placements of the detectors within the LHC and the constraint pertaining to the abort gap, it can be demonstrated trivially that the mathematical solution for this configuration entails a filling scheme that, except the abort gap, is entirely occupied.

In this proposed solution, the optimization of all detectors is emphasized. However, it is important to note that despite 8-fold symmetry in the disposition of the detectors is not valid, because we are considering the real positions, as well as ATLAS and CMS, also LHCb and ALICE will experience an equal number of collisions. This is due to the absence of

any spacing, indeed technically we know that in order to have a specific collision in that detector we should arrange the bunches in the filling scheme in order to “synchronize” with the position of the detectors along the ring, for this reason if we introduce four equidistantly distributed bunches in the same position for the two beams, this will result in a collision occurring at IP1, IP5, and IP2, but not in IP8, because the IP1 and 5 has harmonic 1 and IP2 has harmonic 4, while IP2 has a completely different harmonic. In this case the filling scheme has no harmonic, because there is no spacing between the bunches, it is a constant function, for this reason is synchronized with all positions of the detectors. This ensures that the symmetry of all detectors is maintained, as shown in Figure (3.1).

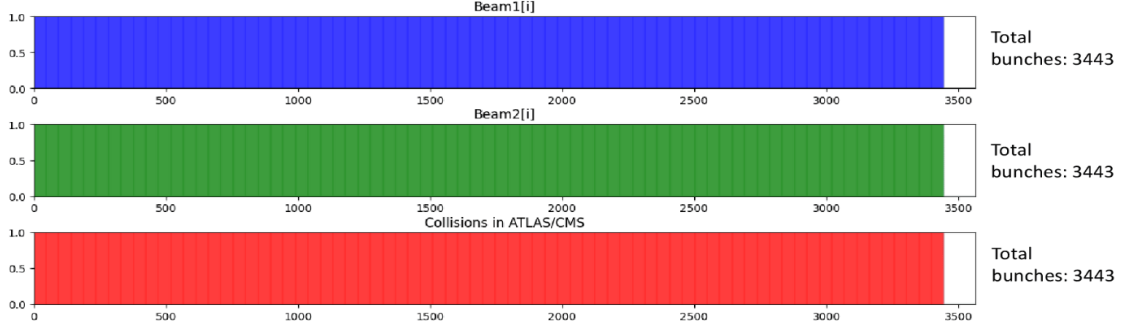
It is crucial to recognize that, despite this equality in the number of collision between LHCb and ALICE, the presence of the abort gap constraint results in a different number of collisions experienced by the other couples of detectors, as shown in Figure (3.1).

3.2.2 2^{nd} case: detectors symmetric, SPS batches

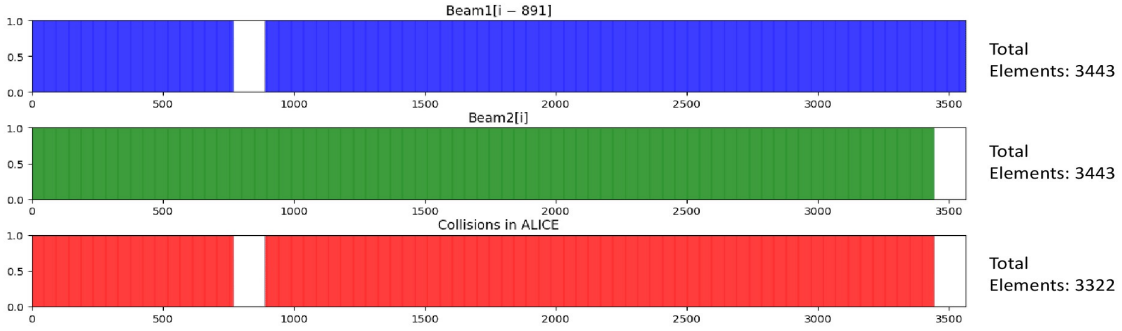
The hypotheses in this particular case differ from the previous one. We will specifically focus on using SPS batches and PS batches for filling the ring, excluding the bunches for tune measurement or INDIVs. The actual spacing between particles will be determined by the RF cavities. Additionally, it is assumed that the 8-fold symmetry is maintained, meaning that LHCb is positioned precisely at the centre of the octave. This results in a perfect symmetry with ALICE in relation to ATLAS. Furthermore, the Abort Gap and the Abort Gap Keeper are no longer taken into account:

$$\begin{aligned}
 & \max_{B1, B2 \in \mathbf{R}^{3564}} \left. \begin{array}{l} (B1 * B2)[0] \\ (B1 * B2)[891] \\ (B1 * B2)[2673] \end{array} \right\} \text{Symmetric detectors' disposition: not realistic, harmonic 4} \\
 & \text{s.t. } \underbrace{\text{SPS}_{\mathbf{n}}(x_1, \text{bun}) = \sum_{i=1}^{\mathbf{n}} \text{PS}(x_i, \text{bun}), \text{ where } x_{i+1} \geq x_i + \text{bun} + 8}_{\text{Machine structure: RF cavity}} \\
 & \underbrace{\sum_{i=1}^{n_{\text{train}}} \text{SPS}_{\mathbf{n}_i}(x_i, \text{bun}) \in B1, B2, \text{ where } x_{i+1} \geq x_i + \mathbf{n}_i \text{bun} + 32}_{\text{Machine structure: RF cavity}}
 \end{aligned}$$

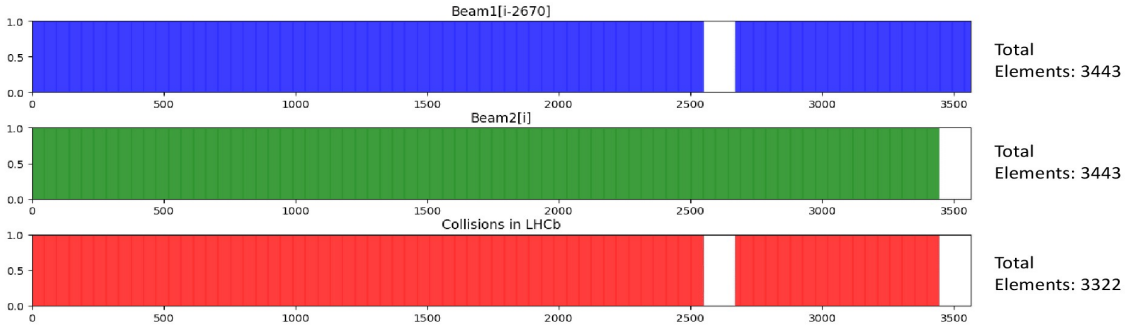
To optimize the number of collisions in the various detectors, it is necessary to strategically fill multiple SPS batches into both beams. This can be achieved by exploiting different degrees of freedom. Assuming that the total number of bunches allocated for filling the ring with both the SPS batches and PS batches has been determined, this particular configuration presents only two degrees of freedom. These correspond to the arrangement of the SPS batches for each beam, with one degree of freedom assigned to each beam. In order to enhance the performance of ATLAS and CMS, it is imperative to consider only one constraint, as depicted in Figure (3.1a). The SPS batches need to be positioned in both beams within the same slot, resulting in an equal number of bunches for each beam. This ensures that the convolution component is maximized, as there is no shift for B1 in relation to these detectors. In this particular scenario, the removal of one degree of freedom in ATLAS/CMS has been introduced to optimize its performance. This ensures that the filling process of one beam uniquely determines the filling of the other beam. Currently, the main goal is to enhance the performance of ALICE and LHCb by optimising the remaining degree of freedom, which pertains to the configuration of SPS batches within a single beam. It is



(a) Shifts for the collisions in ATLAS/CMS



(b) Shifts for the collisions in ALICE



(c) Shifts for the collisions in LHCb

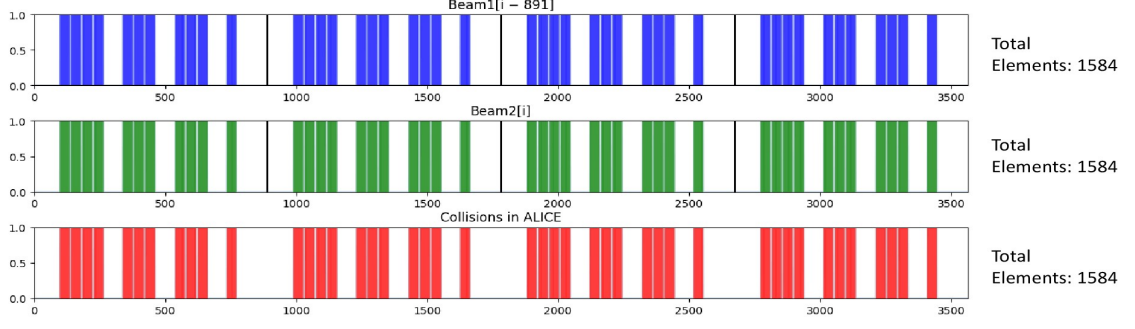
Figure 3.1: Representation of the shifts for Beam1 and consequently the effect of the product for the computation of the number of collisions, for two filling schemes completely filled, except the AG.

worth noting that the detectors' positions exhibit an 8-fold symmetry. However, when considering the shift of Beam1 in the convolution model process described in equation (2.1.2), this symmetry is reduced to a 4-fold symmetry:

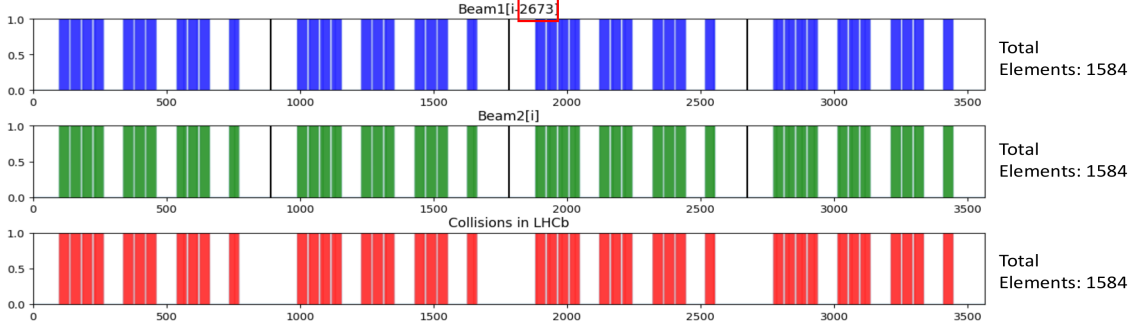
$$\text{mod}(2 * id_{slot}, 3564) = \left[0, \underbrace{891}_{\text{ALICE at } \frac{3564}{4}}, 0, \underbrace{2673}_{\text{LHCb at } \frac{3*3564}{4}} \right]$$

In this particular instance, determining the mathematical layout of the SPS batches is a straightforward task, as the ring consists of a total of 3564 slots. It is necessary for the SPS batches in the filling scheme of B1 to be repeated every 891 slots, which is equivalent to $3564/4$ slots, as illustrated in Figure (3.2).

By neglecting the initial hypothesis that the total number of bunches has been fixed by



(a) Shifts for the collisions in ALICE



(b) Shifts for the collisions in LHCb

Figure 3.2: Representation of the shifts for Beam1 and consequently the effect of the product for the computation of the number of collisions, for two filling schemes created by the repetition of the same structure every quarter of the ring. This is an idealistic case in which the detectors are symmetric

the user and considering the algorithm structure outlined in (3.1), the ultimate objective is to achieve the desired symmetrical configuration while also maximizing the number of injected bunches within the ring, while adhering to the various constraints. One possible approach to attaining the desired arrangement, wherein the SPS batches are injected in the same slots for both beams and a 4 harmonic structure is employed, entails allocating one quarter of the filling pattern for both beams in the identical slots with the $SPS_n(x_i, bun)$, fixed the number of bunches. Afterwards, it is possible to replicate and insert this quarter into the remaining portions, so insuring the preservation of harmonic 4 without the need for manual modification.

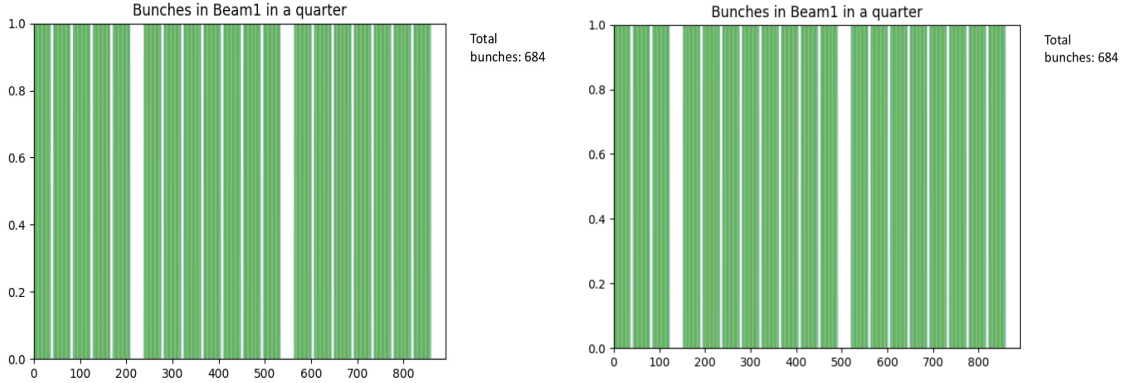
Based on the proposed technique, the only remaining variable is the configuration of SPS batches within one-fourth of the entire ring, with the objective of maximising the overall number of bunches. The aim is to enhance the collision rate in the ATLAS and CMS detectors by maximising the number of bunches within the specified quarter. Based on the information presented in Figure (3.2), it is apparent that this particular strategy would effectively optimise the number of collisions also in ALICE and LHCb. One potential strategy for optimizing the number of bunches, taking into account the maximum number of PS batches inside an SPS batch ($len(PS_{max})$) and the spacing requirements, is to allocate a quarter of the ring to accommodate the longest trains (SPS_{max}) minimizing the empty spaces between the various trains. When it becomes infeasible to insert a train of this type any further, the subsequent step is filling the remaining available space with a shorter train

to complete the ring. This mathematical technique is considered optimal, as seen in Table (3.1).

$\text{len}(SPS_{max})$	# of bunches in the quarter
3	612
4	648
5	648
6	648
7	684
8	684

Table 3.1: Table representing the total number of bunches in a quarter of the ring using the first approach described, using an PS batch of 36 bunches.

Nevertheless, implementing different types of train into a filling scheme might pose significant challenges in practise. In order to mitigate this concern, a limitation has been implemented. Based on this restriction, it is important that the final train to be introduced into the system is not of a length shorter ($\text{len}(SPS_{max}) - 3$). For this, two alternatives for filling a quarter of the ring with the same type of PS batch are illustrated in Figure (3.3). Nevertheless, the first approach given for filling the quarter would result in Figure (3.3b), this configuration is not feasible due to the constraint that has been explained.



(a) Disposition of the PS batches, composed by 36 bunches, in case when $\text{len}(SPS_{max}) = 7PS$. (b) Disposition of the PS batches, composed by 36 bunches, in case when $\text{len}(SPS_{max}) = 8PS$

Figure 3.3: Disposition of the SPS batches composed by different length of PS batches of 36 bunches, in order to fit inside the quarter, following the easier approach

To determine the ideal length, several train lengths have to be tested to determine the extent to which they result in a greater number of bunches. The process in order to test the different trains is exposed here:

1. In this process, we introduce the following notation: n_i represents the length of the SPS (Super Proton Synchrotron) batch, which refers to the number of PS (Proton Synchrotron) batches contained within the SPS batch. The value of n_i ranges from 1 to n , where n is the maximum SPS length specified by the user. M_{n_i} denotes the number of SPS batches with a length of n_i in a given quarter. Additionally, p_{n_i} represents the length of the last SPS batch in the quarter. Finally, bun indicates the number of bunches in each PS batch, which is determined by the user.

2. Based on the given definition: the value of M_{n_i} can be computed as the integer part of $\frac{891}{bun(n_i)+7(n_i-1)+31}$. The value of p_{n_i} can be computed as the integer part of $\frac{(891-31-bun)-(bun(n_i)+7(n_i-1)+31)M_{n_i}}{bun+7}$, where $p_{n_i} \leq n_i - 3$. It is important to note that M_{n_i} is always less than or equal to M_{n_i-1} for all values of n_i greater than 1. Specifically, considering the variable p_{n_i} is for the last train, it is necessary to exclude any elements that cannot be included in the denominator's length with respect to M_{n_i} . This implies that the last PS gap and the SPS gap hasn't to be counted in the total length because they always occur between two batches.
3. The total quantity of bunches in the quarter can be calculated using the formula: $bun[n_i M_{n_i} + p_{n_i} + 1]$.

By performing the aforementioned computation, we acquire a vector that represents the number of bunches within each quarter for various choices of n_i . The highest value from this vector can be selected. Using this n_i and adhering to the other criteria for filling the remaining quarters and positioning the other beam will maximize the number of collisions for all detectors simultaneously.

In this proposed method all detectors are optimized, and as in the previous case, LHCb and ALICE will likewise experience an equal number of collisions alongside ATLAS and CMS.

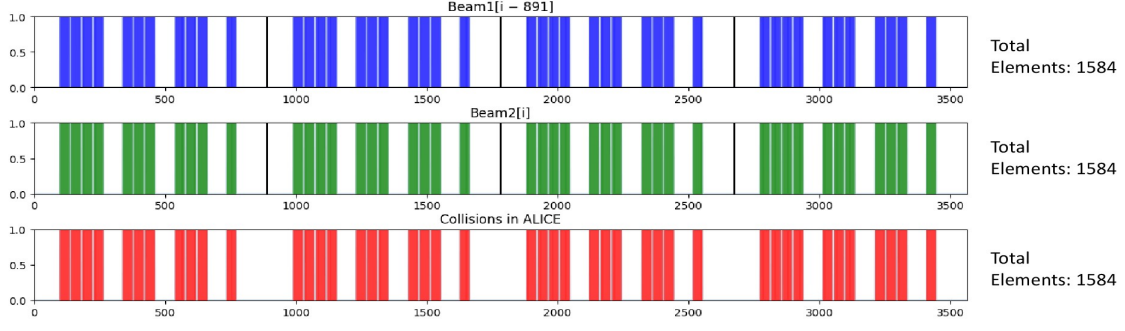
3.3 Realistic cases

To address the practical scenario, it becomes imperative to delete all assumptions made in the previous chapter's analyses. Unfortunately, it is no longer feasible, mathematically speaking, to accomplish the original objective of determining the optimal filling scheme that maximizes the performance of all four detectors simultaneously. The reasons for this limitation would be elucidated.

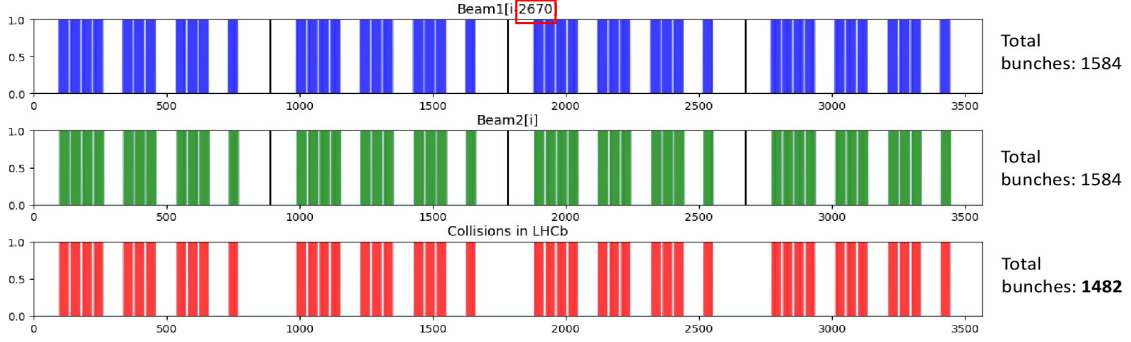
In practical scenarios, it is seen that the filling scheme exhibits empty gaps between batches due to the presence of RF cavities in the accelerator chain. Furthermore, the positioning of the detectors does not adhere to an 8-fold symmetry. Given the requirement for equal numbers and types of trains on both beams, the positioning of the trains, using the two degrees of freedom exposed in the previous section, is fundamental in order to achieve our goal. The methods described in the previous chapter allow for using the symmetry between ATLAS and CMS, thereby maximizing their potential. This can be achieved by solely maintaining a consistent spatial distance between the beams for each SPS batch. Notably, this distance is set to zero. However, compared to the other cases, now the harmonic 8 is no longer adhered to, thereby it's needed to use also the distance between the SPS batches within the same beam. This enables the maximization of one of the two remaining detectors, due to the absence of symmetry between them. Indeed, in Figures (3.4) and (3.5) two distinct filling schemes are depicted, both utilising the same SPS batches but arranged in different manners. It is evident that while the filling scheme in Figure (3.4) maximises the number of collisions in ALICE, this does not necessarily imply that LHCb is also maximised. Actually, an alternative configuration, as shown in Figure (3.5), presents that the number of collisions in ALICE is reduced while the number of collisions in LHCb is increased.

In summary, it may be stated that it is not feasible to identify a filling scheme that maximizes all four detectors simultaneously. However, it is rather straightforward, according to the approach presented in the preceding section, to identify a filling scheme that maximizes three out of the four detectors.

However, in the preceding chapter it was not solely assumed that the detectors were symmetric; indeed, three other hypotheses have been presented:



(a) Shifts for the collisions in ALICE



(b) Shifts for the collisions in LHCb

Figure 3.4: Representation of the shifts for Beam1 and consequently the effect of the product for the computation of the number of collisions, for two filling schemes created by the repetition of the same structure every quarter of the ring. This is a real case in which LHCb is disposed in the right position

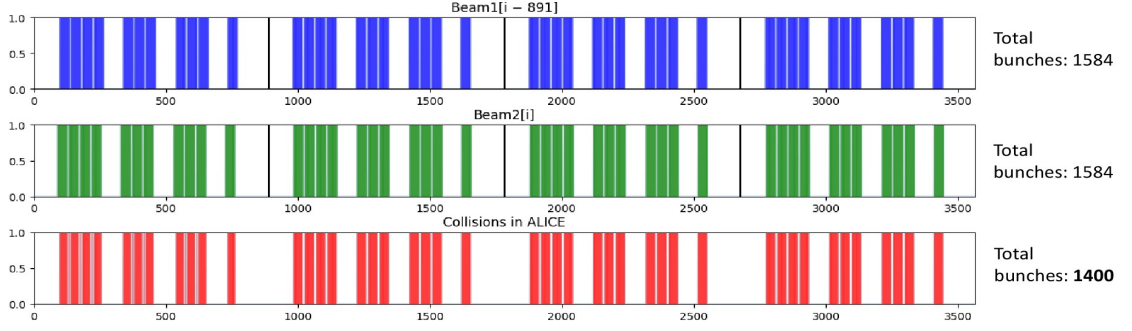
- The terms “Abort Gap” and “Abort Gap Keeper” were not taken into consideration.
- The selection of bunches for the tune measurement was not taken into consideration.
- The opportunity to incorporate INDIV into the filling system was not available.

If these hypotheses are eliminated, the complexity of the problem will increase. Specifically, the inclusion of these limitations in the optimization problem will no longer yield the desired outcome of maximizing three detectors simultaneously using a harmonic analysis. Why?

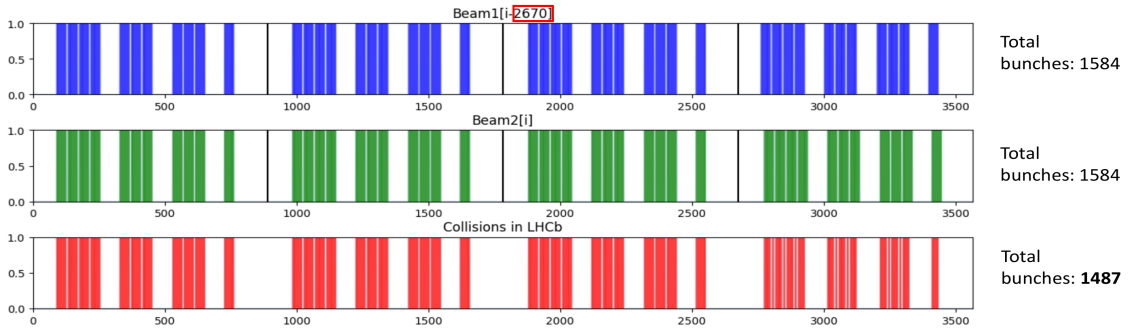
Because each of these constraints damages the harmonic structure behind the analysis of the second degree of freedom. Taking into account the Abort Gap, it can be concluded that the entire ring is no longer available for injecting bunches, resulting in a reduction of 121 slots in its length. Consequently, it is no longer feasible to replicate all the SPS batches within the same beam over a distance of 891 slots due to insufficient space. As a result, the previous studies are rendered invalid.

Based on the aforementioned line of reasoning, it is straightforward to further explore the question at hand. Given the many constraints inherent in the LHC, the question is to figure out the possibility of identifying a filling scheme that optimizes the simultaneous maximization of three out of the four detectors.

Before addressing these questions, it is important to clarify a detail regarding the constraints. The most intricate constraint that needs to be taken into account is the final one



(a) Shifts for the collisions in ALICE



(b) Shifts for the collisions in LHCb

Figure 3.5: Representation of the shifts for Beam1 and consequently the effect of the product for the computation of the number of collisions, for two filling schemes created with a particular approach. This is a real case in which LHCb is disposed in the right position

(the INDIVs), as it is observed to be the one that most violates the harmonic analysis. Therefore, the optimization problem will initially be solved without considering INDIV, and subsequently, the complete problem will be addressed.

3.3.1 3rd case: Real filling scheme without INDIV

From now on, the mathematical methodology employed in the preceding chapters has been abandoned. The algorithm presented in this chapter adheres to the same conceptual framework as the previous chapters, although lacking formal mathematical proofs. Indeed, these algorithms have solely undergone numerical testing.

This algorithm is specifically designed to address the optimization problem with all the possible constraints, apart:

$$\begin{aligned}
 & \max_{B1, B2 \in \mathbf{R}^{3564}} \left. \begin{array}{l} (B1 * B2)[0] \\ (B1 * B2)[2670] \end{array} \right\} \text{Real detectors' disposition} \\
 & \text{s.t. } \underbrace{\exists! \text{ PS}(x, 12) \in B1, B2, \text{ where } x < 36}_{\text{Machine diagnostic: bunches for tune measurement}} \\
 & \underbrace{|x_2 - x_1| \geq \text{bun}_q, \text{ for } \text{PS}(x_1, \text{bun}_q)_{B1}, \text{PS}(x_2, \text{bun}_q)_{B2}}_{\text{Machine diagnostic: bunches for tune measurement}} \\
 & \underbrace{\text{SPS}_{\mathbf{n}}(x_1, \text{bun}) = \sum_{i=1}^{\mathbf{n}} \text{PS}(x_i, \text{bun}), \text{ where } x_{i+1} \geq x_i + \text{bun} + 8}_{\text{Machine structure: RF cavity}} \\
 & \underbrace{\sum_{i=1}^{n_{\text{train}}} \text{SPS}_{\mathbf{n}_i}(x_i, \text{bun}) \in B1, B2, \text{ where } x_{i+1} \geq x_i + \mathbf{n}_i \text{bun} + 32}_{\text{Machine structure: RF cavity}} \\
 & \underbrace{\delta(x), \text{SPS}_{\mathbf{n}}(x, \text{bun}), \text{ where } x \leq \text{AG} = (3564 - 121)}_{\text{Machine protection: Abort Gap and Abort Gap Keeper}}
 \end{aligned}$$

and only for protons studies in LHC.

Indeed the LHC employs a hierarchical approach in selecting detectors for various analyses and studies. Specifically, for proton studies, the preferred detectors are ATLAS and CMS, followed by LHCb, with ALICE being the least prioritized. Conversely, for ions studies, ALICE is the preferred detector, followed by ATLAS and CMS, while LHCb is the least prioritized. Hence, the algorithm that will be presented in this study aims to prioritize the maximization of ATLAS and CMS, followed by LHCb. However, it is worth noting that the approach may be readily expanded to encompass the maximization of ATLAS, CMS, and ALICE.

One final aspect to emphasize is that this algorithm was discovered following an extensive period of data analysis, which involved studying the filling schemes employed in the Large Hadron Collider over previous years. Additionally, the investigation examined the underlying structure of the “Filling Scheme Editor” used by the LHC Programme Coordinator (LPC) ([LPC]) and its “Pre-fill” algorithm. The objective was to identify a pattern or rule that would optimize the filling process.

Optimized algorithm Pre-fill without INDIV

The underlying concept of this method is to incorporate the harmonic structure previously discussed in the preceding chapter, while addressing the increased complexity of the current scenario and trying to expand upon that approach. We have taken into account all the constraints, except the INDIV constraint. Although the other constraints we have introduced may disrupt the symmetry of the problem, we are fortunate that these two constraints are satisfied at the beginning and end of the ring. From a mathematical perspective, this violates the principles of harmonic analysis. However, by employing certain strategies, we can adapt the previous analysis to accommodate this particular scenario.

The full procedure behind the algorithm is presented in Appendix A.2 in which is presented in all the passages, considering the different cases.

Given that the present analysis involves a reevaluation of a harmonic approach presented before, it is important to note that the mathematical proof for these passages is lacking. However, to address this limitation, we conducted a numerical assessment of the algorithm using a Monte Carlo simulation. The subsequent section will discuss into the details of this simulation, which plays a central role in our investigation.

3.3.2 4th case: Real filling scheme with INDIV

Taking into account all the limitations, including those pertaining to the individual particles (INDIVs), the problem's structure would become significantly intricate, as represented in (3.1). This is mostly due to the arbitrary nature of the INDIVs' injection, with the single restriction of head-on collisions between them in specific detectors. Due to the aforementioned factors, it is feasible to insert these entities into any available positions inside the ring. Given that the harmonic analyses conducted in the preceding chapter are deemed irrelevant for this particular scenario, it becomes very challenging to devise an algorithm that effectively arranges the SPS batches and INDIVs according to the user's specifications while simultaneously maximizing the number of collisions. The issue in this particular case does not lie in the optimization of ATLAS/CMS, as it can be easily achieved by placing the batches for both beams in the same slot. However, the challenge arises when attempting to maximize LHCb, as the conventional copy and paste method can no longer be employed due to the potential presence of an INDIV requested by the user in the middle of a quarter.

Due to this reason, we have chosen to undertake a comprehensive examination of the optimization problem via an alternative methodology. This approach diverges from the previously described harmonic and theoretical methods, instead adopting a pragmatic perspective by employing a Monte Carlo simulation.

Monte Carlo simulations using a three maker

The conceptual framework underlying this approach diverges significantly from the other approaches discussed. Unlike the algorithmic methods, this simulation does not generate a filling scheme based on input data from the user about the filling scheme. Rather, it operates on a pre-existing filling scheme. In what manner can this be accomplished? The algorithm calculates the vacant intervals between the SPS batches of the filling scheme that exceed the size specified by the RF cavities (800 ns), as shown in Figure (3.6).

These intervals are then stored in a vector. From this computation, it is feasible to determine the SPS batches that can be shifted and the direction in which they can be shifted. A random element, distinct from zero, is selected from the vector and the corresponding SPS batch is shifted by one slot in that specific direction for both beams simultaneously. This process generates an alternative filling scheme that adheres to the constraints of the optimization problem. The SPS batches has been shifted in both beams in order to maintain the number of collisions in ATLAS/CMS unaffected, as it is already maximized. However, by altering the spacing between the SPS batches within the same beam, alternative configurations of the filling scheme can be explored. These configurations may result in a greater number of collisions in either one or both of the other two detectors. Furthermore, to maintain the integrity of the constraints regarding the movement of INDIVs and to prevent the occurrence of INDIV-INDIV collisions at an interaction point different from ATLAS/CMS, the INDIVs involved in collisions are simultaneously displaced to ensure a constant separation distance, in order to maintain the collision occurring within that specific detector.

Ultimately, by repeatedly applying this arbitrary transformation and preserving each intermediate iteration, it becomes feasible to generate a Monte Carlo simulation that exhibits

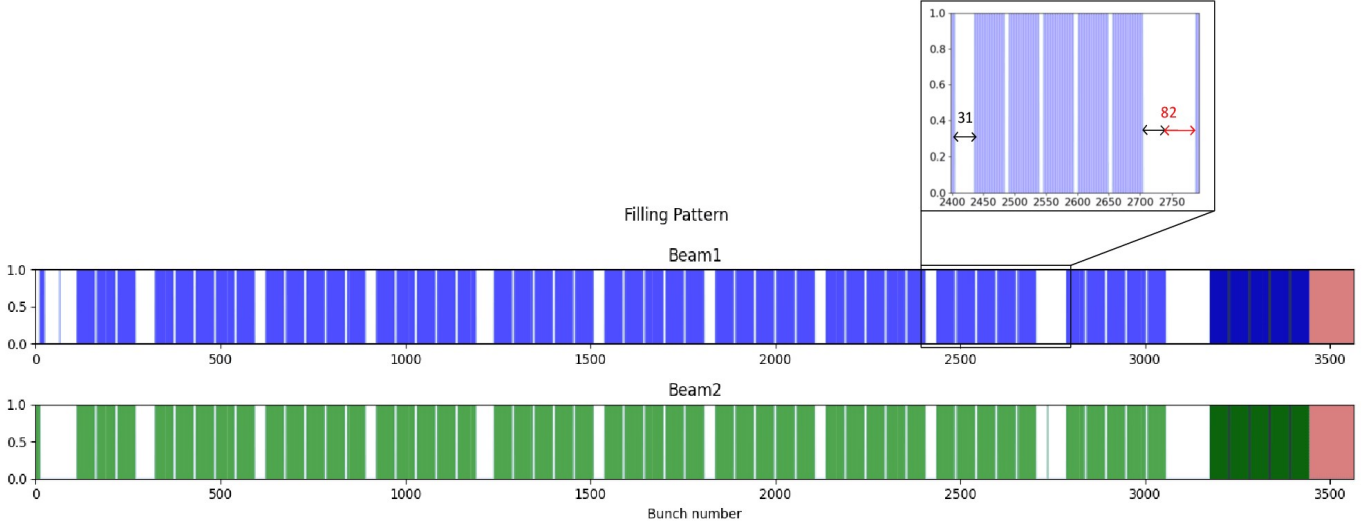


Figure 3.6: Example of filling scheme used at CERN, B1 in blue and B2 in green, same as shown in Figure (2.3). In the zoom, it is represented the distance between SPS batches, underlining the gaps longer than 800 ns (31 slots).

alternative arrangements for the filling scheme that has different numbers of collisions for ALICE and LHCb experiments.

From a computational perspective, the Monte Carlo simulation in question does not impose a significant burden in terms of processing cost. This is mostly due to the fact that the simulation primarily consists of sub-vector shifts. The main issue is the computational cost of calculating the number of collisions in ALICE and LHCb at each intermediate step, given that the number of iterations could be as high as $\sim 50000/100000$. For this reason, the first chapter focused on the computational cost of the optimized script for computing the number of collisions.

An additional noteworthy aspect of this simulation, from an informatics perspective, pertains to memory allocation. Given the substantial number of iterations involved and considering that it is necessary to store not only the varying collision counts for LHCb and ALICE but also all the filling schemes. Those are crucial for assessing compliance with constraints and facilitating the continuation of the simulation from intermediate stages. Consequently, meticulous attention must be paid to the memory storage of this simulation:

- One possible approach to efficiently save the intermediate step is by using `numpy.array` to store the two boolean vectors representing the filling scheme of the two beams. Given that the memory allocation for a boolean number is typically 1 byte, and assuming a vector length of 3564 and a potential number of steps of approximately 50000/100000, the required memory storage can be calculated. Approximately 713 megabytes.
- One trivial improvement can be obtained by leveraging the fact that the two beams, except for the initial bunches employed in the tuning measurement, are situated in identical slots. Instead of storing two separate boolean vectors for each beam, it is feasible to save a single boolean vector that represents one beam and includes the injection slot for the tune measurement of the other beam. This approach reduces the required memory storage. Approximately 357 megabytes.
- An intriguing potential enhancement might involve leveraging a specific function from the numpy library, namely `numpy.packbits` ([npbits]). This function facilitates the

conversion of components within a binary-valued array into bits stored in an uint8 array. As a result, the memory allocation of the array can be reduced by a factor of up to eight. Due to this rationale, employing this methodology results in a further reduction in memory capacity. Approximately 56 megabytes.

- The approach employed in the final script differs as well. By leveraging the fact that the various SPS batches, excluding the bunches used for tune measurements and the INDIVs, consist of a series of PS batches with an equal number of bunches, it becomes feasible to optimize storage by retaining only the initial and final slots of each SPS batch. By using this information, along with the knowledge that the PS empty gap covers 7 slots, it becomes possible to accurately reconstruct all SPS batches uniquely. Given that the position slots are represented as integers and the ring has a length of 3564 slots, it is typically observed that there are no more than 20 injections in each beam. By employing this methodology and storing the positions of the slots as unsigned 16-bit integers, it is feasible to attain a memory storage of around 8 megabytes.

The simulation aims to explore a wide range of potential configurations derived from a predetermined filling scheme, with the objective of identifying the optimal configuration among the numerous possibilities. To enhance the exploration of various configurations, it is beneficial to conduct multiple parallel Monte Carlo simulations starting from a common filling scheme. This approach allows for the examination of a wider range of potential configurations. Additionally, when aiming to maximize the number of collisions in the LHCb and ALICE detectors, it is more advantageous to select the most optimal filling scheme in terms of collisions for these remaining two detectors. Subsequently, this selected scheme can be employed as a starting point for subsequent simulations, thereby investigating the potential for further improvements in collision outcomes.

Based on the aforementioned objectives, it has been determined that the implementation of a Monte Carlo simulation using a three-maker approach is appropriate. This approach involves generating a set of filling schemes, each of which produces different descendants. During the transition from one generation to the next, only the filling schemes that result in a greater number of collisions in ALICE at each value of number of collision in LHCb (i.e., the most optimal ones) are selected. Subsequently, the next generation is regenerated based on these selected filling schemes, and the process continues iteratively.

Based on the pragmatic and random nature of the Monte Carlo simulation, it is not possible to definitively claim that this simulation will yield the optimal solution. This assertion is further weakened by the structure of the three maker, as the transition between consecutive generations is comparable to the algorithm of the steepest descent method, which is known to potentially converge to local minima.

Results of the MonteCarlo simulations

The objective of this Monte Carlo simulation is to identify a more optimal configuration, with respect to the number of collisions in LHCb and ALICE, based on the given constraints of the optimization problem. Hence, this simulation, which investigates several potential filling schemes, could serve as a valuable tool to numerically test the methodology, results, and assertions presented in the preceding sections.

- As an initial step, we can conduct a test on the simplest scenario, which involves symmetric detectors. This test will use the structure of the SPS batches, namely the 2nd case discussed in the preceding section. We will begin the Monte Carlo simulation with the solution that we have previously determined to be the most optimized.

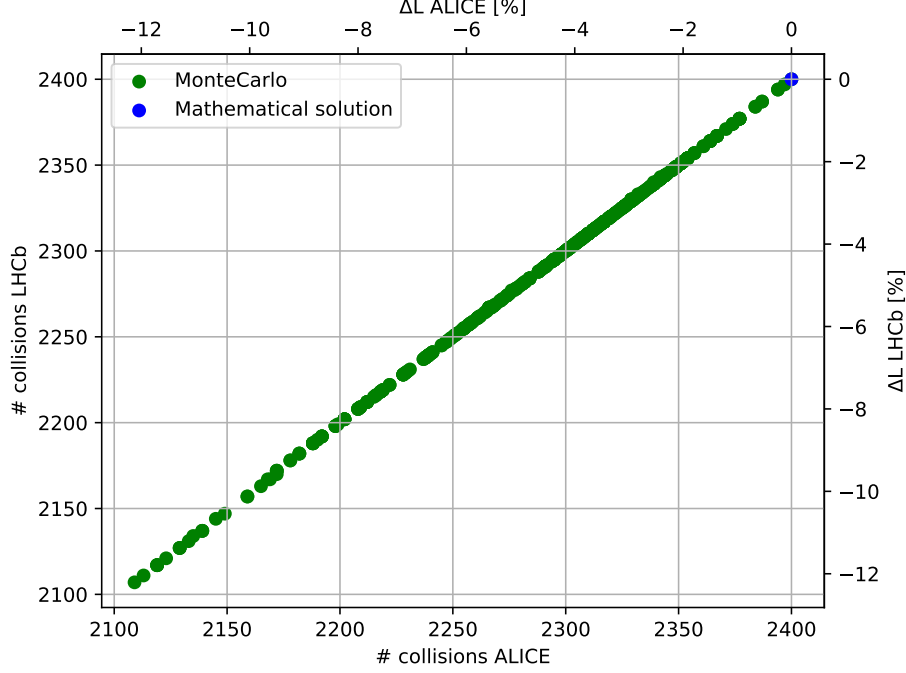


Figure 3.7: Different values of number of collisions for LHCb and ALICE in all the configurations reached by the Monte Carlo, considering LHCb shifted perfectly symmetric to ALICE, starting from a used filling scheme in the LHC with SPS batches composed by 7 or 5 PS batches of 36 bunches. Starting the Monte Carlo simulation from the optimized solution found in the previous section.

The Figure (3.7) demonstrates that the Monte Carlo results align with the findings presented in the corresponding section. Specifically, the number of collisions observed in ALICE and LHCb exhibits perfect equality. Furthermore, all the data points conform to the diagonal line on the plane, indicating that the simulation does not identify a more optimal configuration compared to the theoretically obtained.

- For the second test, a more realistic scenario may be examined by evaluating an actual filling scheme that incorporates several SPS batches. All restrictions, except for the INDIVs constraint, will be taken into account. This corresponds to the 3rd case that was discussed. In the preceding section, a mathematical proof establishing the optimality of the algorithm in terms of maximizing collisions in LHCb or ALICE was not obtained. However, through the use of Monte Carlo simulations, various configurations were numerically reached. The objective was to evaluate the algorithm's performance and determine if the outcomes align with the anticipated promising results. Consequently, the simulation was started from the filling scheme determined by the algorithm.

The simulation depicted in the Figure (3.8) provides confirmation for the assertions made in the preceding section. Specifically, approaching a more realistic scenario, it becomes increasingly challenging to identify a single optimal solution that maximizes all detectors simultaneously. This is due to the possibility of encountering a filling scheme that results in a greater number of collisions in one detector while adversely affecting the other. Furthermore, the Monte Carlo simulation successfully identified more optimal configurations in terms of collision count in the ALICE experiment, but not in the LHCb experiment, which was maximized by the algorithm. This can be seen as a numerical validation of the algorithm's effectiveness.

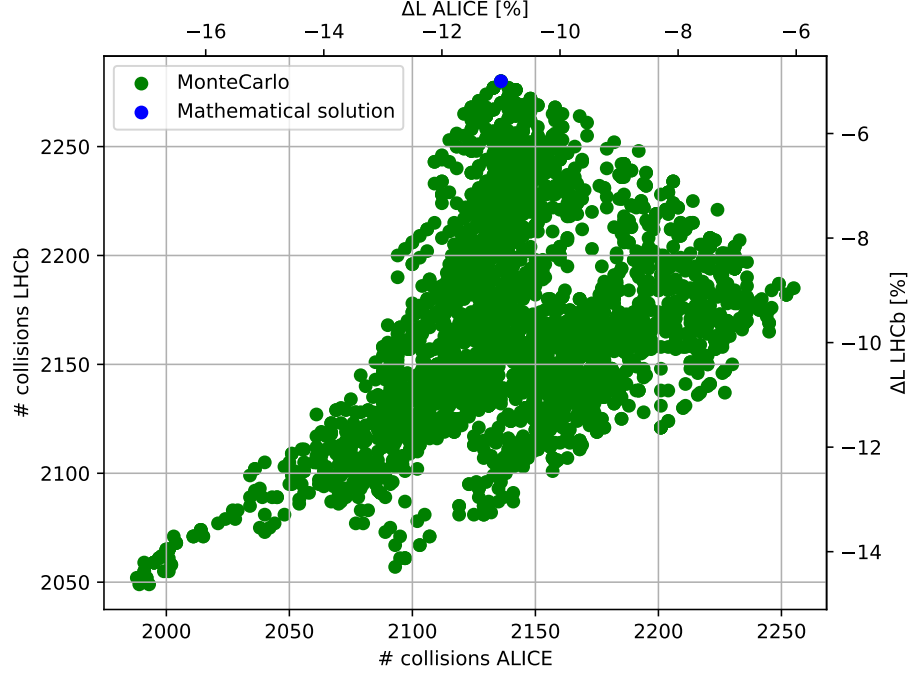


Figure 3.8: Different values of number of collisions for LHCb and ALICE in all the configurations reached by the Monte Carlo, considering LHCb and ALICE in the right positions, starting from a used filling scheme in the LHC with SPS batches composed by 7 or 5 PS batches of 36 bunches without any INDIV. Starting the Monte Carlo simulation from the optimized solution found with the algorithm for the maximization of the number of collisions in LHCb in the previous subsection.

The implementation of this Monte Carlo simulation in a practical filling scheme with INDIV has yielded the most intriguing outcome. Specifically, we have examined the most common filling scheme employed during the RUN 3 in LHC in 2022. The selected filling scheme used in the LHC was chosen from a comprehensive dataset. The simulation benefited the Three Maker structure, as described earlier, employing three generations. The process involved an initial longer step consisting of 80,000 shifts, followed by subsequent generations with shorter steps of 3,000 shifts. This approach facilitated the exploration of various configurations in a more expedited manner. Our predicted outcome would have been that the filling scheme employed in the LHC would align with the boundaries of the Monte Carlo simulation. Consequently, we would not have been able to identify a superior filling scheme in terms of collisions within the ALICE and LHCb detectors. However, we could have identified alternative filling schemes that would favour one detector over the other.

However, the outcome depicted in Figure (3.9) contradicted our first beliefs. The simulation successfully identified different configurations for the arrangement of SPS batches that had the potential to increase the amount of collisions in both detectors concurrently. Hence, it can be concluded that alternative filling schemes proved to be more efficient compared to the scheme employed in the previous year in the LHC. In particular, this approach has resulted in an increase of 2% in the number of collisions for each detector at maximum. Although this outcome may not be considered a pivotal finding or so overwhelming, it is entirely cost-free, as it merely serves as a means to maximize our available resources and expertise. Furthermore, we cannot guarantee that we have identified the most optimal boundary, as indicated by the uncertainty related to our ability to attain the global optimum through this simulation.

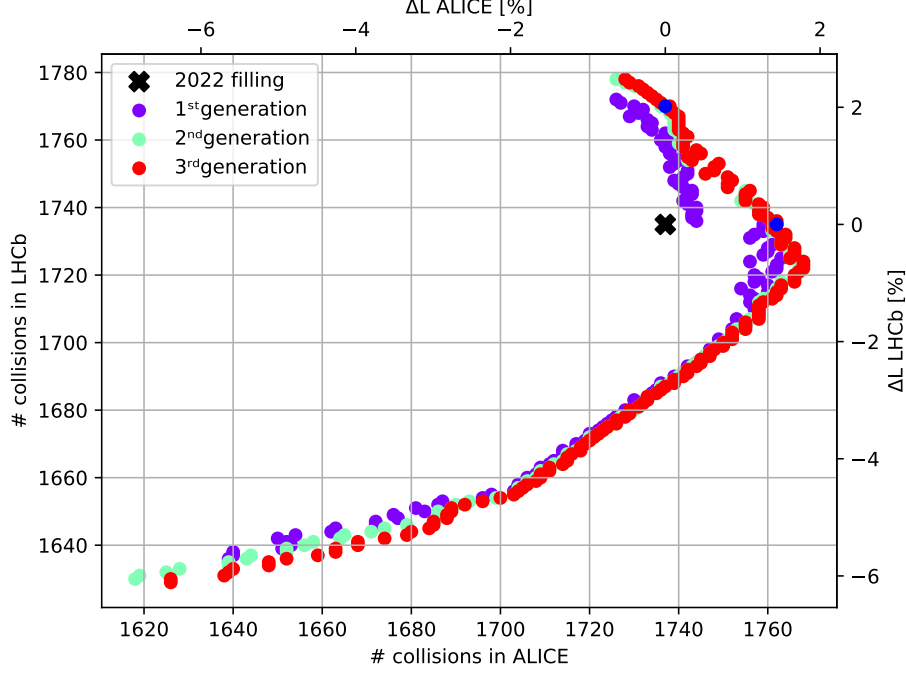


Figure 3.9: Different values of number of collisions for LHCb and ALICE in all the configurations reached by the Monte Carlo, considering LHCb and ALICE in the right positions, starting from the most used filling scheme in 2022 in the LHC with SPS batches composed by 5 or 3 PS batches of 36 bunches and two INDIV each beam. Starting the Monte Carlo simulation from the filling scheme used in the LHC, represented in the plot by a black cross (2461b_2448_1737_1735, notation NumberOf-Bunches_ATLAS/CMS_ALICE_LHCb), and the two blue points are two configurations that fit on the boundary of the Monte Carlo simulations and have the same number of collisions one in LHCb and the other in ALICE of the initial configuration.

3.4 Conclusions and following studies

The objective of this analysis was to employ optimization and algebraic techniques to capitalize on the problem's symmetry and constraints, outlined using the model defined in the Section (2.1.2), with the aim of maximizing the number of collisions among designated detectors. A threshold has been established, beyond which the identification of an optimal solution in analytical terms becomes non-trivial. Consequently, a more pragmatic shortcut has been devised to overcome the limitations of our analytical capabilities, by studying the problem using a brute force approach, exploiting the python optimization presented in the Section (2.2). This approach has yielded noteworthy results and improvements to a filling scheme that was implemented in the LHC. Moreover, this brute force strategy has been employed for some experimental session, namely Machine Development, in 2023.

An extension of the aforementioned approach, which has been already incorporated in the code, can be to use it not only for the analysis of protons filling scheme in the LHC, but also for heavy ions. In this particular scenario, the problem's structure undergoes modifications, affecting both the objectives and the limitations. Specifically, in relation to the objective, there is no longer a necessity to prioritize the maximization of ATLAS/CMS, followed by LHCb, and finally ALICE. The hierarchical structure in this scenario has been completely reversed. The objective is to optimize the performance of ALICE, ATLAS, CMS, and lastly

LHCb, with the understanding that the performance of LHCb may potentially be lower as that of the other experiments but at least larger than zero. Hence, it is no longer necessary for the batches of the two beams to be injected into the same slots; rather, they might be asymmetric, properly because maximizing the number of collisions in ATLAS/CMS is not the first goal. The only need is that the injection remains consistent, ensuring an equal total number of bunches between the two beams. Hence, the Monte Carlo simulation has the capability to iteratively change the SPS batch positions of the two beams in diverse manners.

Additionally, there is a slight modification in the limitations. In the case of heavy ions, the intervals between bunches cannot be solely 25 ns (equivalent to 7.5 m). Instead, the distance is required to be doubled at minimum. In the filling schemes employed for heavy ions in recent years, the intervals between bunches were adjusted to 50 ns, 75 ns, and 100 ns. Consequently, this alteration in the spacing between bunches also affects the spacing between PS batches, which varies depending on the specific scenario. For further information regarding this extension, interested readers may refer to the work of ([Co16]).

The aforementioned expansion was used during the operational phase of ions in 2023 to delineate the optimal filling scheme options with respect to the constraints of the given scenario.

An additional expansion, drawing inspiration from the work ([Fa21]), involves doing an analysis of the aforementioned problem within the context of hybrid protons filling scheme. In the present year, a novel filling scheme has been implemented for the first time in the LHC during the running period. This scheme is characterized by the utilization of two distinct types of PS batches. The objective of this filling technique is to mitigate the adverse phenomenon known as electron-cloud ([Zi02]).

Furthermore, an additional intriguing improvement can be implemented in the Monte Carlo simulation. In particular, can be defined a useful metrics in the “collisions plane”. After conducting the initial iteration of the Monte Carlo simulation, the selection process will identify not only the most optimal filling schemes but also the one more distant from the initial one. Following this, the subsequent generation is regenerated using the aforementioned selected filling schemes, and this iterative process persists. This alternative methodology exhibits a higher degree of randomness compared to the previous one, and it slightly deviates from the steepest descent method.

Finally, it would be intriguing to expand the scope of this study by using the Monte Carlo simulation to not only include the shifting of SPS batches in the filling scheme, but also the PS batches considering all the relevant constraints of the given scenario.

Chapter 4

Mathematical model: luminosity from emittance

The objective of this thesis is to employ various mathematical techniques to examine the calculation of luminosity, aiming not only to maximize it but also to take advantage of this value's potential in deriving significant metrics that are essential for characterizing the efficiency of the accelerator.

The luminosity formula presented in the Introduction (Eq. 1.1) comprises two primary components: the number of collisions in each interaction point (n_{IP}), which was the focus of analysis in the initial phase of this project aimed at maximizing it in various filling scheme configurations, and the quadruple integral ($\mathcal{L}_{col}$), which will be the subject of investigation in the subsequent part of this thesis.

The main goal of this part of the project is to derive the emittance, by the “inversion” of the quadruple integral.

4.1 The beam emittance and its importance

In the field of accelerator physics, as seen in the Introduction 1.1 and in Chapter 2, a beam consists of a collection of particles, moving all with the same total energy. In the LHC, the beam within the pipe is not continuous, it is organized into clusters known as bunches. These bunches consist of a collection of particles, typically protons, densely packed together, and these bunches are spaced by ≈ 7.5 m. The number of particles contained within each bunch is approximately $1e11$. The beams, in a cycle, run a distance of ≈ 27 km, and they are designed to collide with each other as many as possible. The engineering challenge of this goal is particularly clear, considering that the transverse area of each beam is on the order of 1×10^{-6} m, basically it's like we would like to make collide two needles throwing them from Geneva to Chicago.

To determine the number of collisions, as discussed in Chapter (2), we employed a model that treated the bunches as discrete entities, disregarding the density profile. Instead, in order to determine the luminosity, it is imperative to acknowledge that these bunches exhibit varying density distributions in the three spatial directions, that represent the distributions of the bunch particles, commonly considered Gaussian. Specifically. The emittance is closely linked to the variance of these distributions, and this also explains its significant correlation with the luminosity.

It is evident from this short introduction why it is important to measure this beam parameter, and it will be expressed in the section how this correlation can be expressed from a mathematical point of view.

4.1.1 The definition of beam emittance

The emittance of a single particle as well as the emittance of a beam can be defined.

Courant-Snyder invariant or Single particle emittance

Using a classical mechanical approach, the goal, given a phase space trajectory, it is to understand the invariant of the motion.

To establish the concept of emittance, it is necessary to formulate the equations of motion using Hamiltonian formalism:

$$\begin{cases} \dot{i} = \frac{\partial H}{\partial p_i} \\ \dot{p}_i = -\frac{\partial H}{\partial i} \end{cases} \quad i = x, y, z, \quad (4.1)$$

where H is the Hamiltonian of the system. If it does not depend explicitly on time, as shown ([dgaranin]), it represents the total energy of the system that is invariant; in this case, there is only one phase space trajectory originating from any point in the phase space, and distinct trajectories cannot cross. In this particular analysis, we examine the case where the motion along each axis is separated from each other, the motion along the axes are all decoupled. The idea is to find other factor function of x and p_x that is commutative with respect to the Poisson brackets.

In particular, in a beam line where the longitudinal magnetic field is considered to be negligible, the potential vector $\vec{\Phi}$ of the magnetic field is perpendicular to the magnetic field itself, indicating a longitudinal orientation. In the present scenario, the transverse constituents of the particle momentum, denoted as $\vec{p} - e\vec{\Phi}$, are equal to the transverse component of the conjugate momentum, represented as $\vec{p}$. Consequently, the derivative on time x' is linked to p_x through a linear relation:

$$p_x = m\beta_r\gamma_r c x'$$

where m is the rest mass of the particles, the relativistic $\beta_r = \frac{v}{c}$ and $\gamma_r = \frac{1}{\sqrt{1-\beta^2}}$ are related to the particle velocity and so to the particle's energy. Hence, in a ring without acceleration and devoid of longitudinal magnetic field, the particle's energy remains constant, thereby preserving the ratio between x' and p_x . In this case, it is possible to define another invariant of the motion in terms of x and x' , that is called the Courant-Snyder invariant (also referred as single particle emittance):

$$\epsilon_{CS,x} = \left(\frac{\gamma}{2} x^2 + \alpha x x' + \frac{\beta}{2} x'^2 \right). \quad (4.2)$$

where:

- $\sqrt{\beta}$ represents the particle envelope per unit emittance,
- $\sqrt{\gamma}$ represents the particle divergence per unit emittance.
- $\sqrt{\alpha}$ is proportional to the correlation between x and x' .

these parameters, as it will be explained later are defined as optics function or Twiss parameters, each of them has a geometric interpretation as shown in Figure (4.1).

Specifically, the particle within the transverse plane undergoes oscillations along the accelerator, as represented in Figure (4.1) on the left. It has been demonstrated in ([Se92]) that an individual particle, in a linear accelerator, traces an elliptical path in the transverse

phase space, as represented in Figure (4.1) on the right. The area contained in this trajectory in the phase space is equal, by a factor 2π , to the transverse emittance ϵ of a single particle defined in Equation 4.2. It is possible to show that this definition is related to the Hamiltonian mechanics, indeed it is possible to show that the Hamiltonian action from 4.1 is perfectly equal to the single particle emittance. In alternative phrasing, the concept of single particle emittance quantifies the magnitude to which a particle's trajectory diverges from an ideal reference trajectory during its passage through the ring.

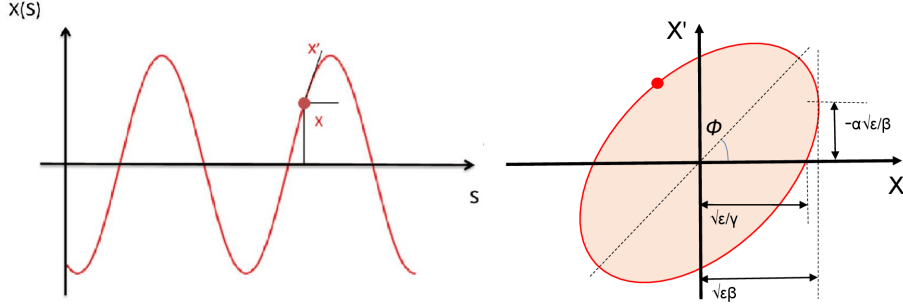


Figure 4.1: On the left of the picture is represented the particle's trajectory with respect to the reference one. It is worth noting its oscillating behaviour. On the right, its movement in the phase space is represented. It follows an ellipse shape. The geometrical meaning of α, β and γ parameters (the optics functions) is also depicted.

In a circular accelerator such as the LHC, it can be observed that the emittance is not constant during the acceleration ramp. However, the energy normalized emittance

$$\epsilon_{CS,xN} = \beta_r \gamma_r \epsilon_x$$

is an invariant in a linear machine. The reason for this phenomenon is that while the product of β_r and γ_r increases in direct proportion to the momentum of the particle ($\frac{p_x}{mc} = \beta_r \gamma_r$), the emittance is proportional to the inverse of the momentum.

In the following Section, we will generalize the definition made for a single particle in the case of an ensemble of particles, trying to keep its main characteristic, the invariance. For a more comprehensive understanding of this definition and its generalization for an ensemble, interested readers can refer to the works of ([Bu94]) and ([Hi21]) for further details.

Statistical beam emittance

The beam can be effectively modelled by the Hamiltonian system, wherein the system's N particles can be seen as a Hamiltonian system with $3N$ degrees of freedom, not including the internal degrees of freedom of the particles. Consequently, the canonical phase space should be considered to have dimensions of $6N$, accounting for the momenta as well.

In this project, the model is simplified by assuming that the N particles are indistinguishable and do not interact with each other. Consequently, only the phase space of a single particle is considered, which serves as a representative for the entire system. The $6N$ phase space describes the state of the beam at a specific time t . This state is represented by a collection of N points, denoted as $P_j(t)$, with each point corresponding to an individual particle.

In this particular investigation, we examine a scenario where the longitudinal motion along the beam axis is fully decoupled from the motion in the transverse plane of the beam axis. Consequently, the 6-dimensional phase space can be divided into a longitudinal phase space consisting of 2 dimensions and a transverse phase space consisting of 4

dimensions. Furthermore, assuming uncouple transverse motion, the transverse phase space can be divided into two 2D phase spaces. Eventually, the 6-dimensional phase space undergoes a reduction and simplification process, resulting in the formation of three distinct 2-dimensional phase spaces, each corresponding to a certain spatial component.

In the scenario of ignoring mutual interactions and coupling among the three coordinates of a particle, it becomes feasible to establish the beam emittance for each individual degree of freedom, that is invariant for all the phase spaces.

In this section of the thesis, the focus will be on the transverse emittance, so analysing only those in the horizontal and vertical planes. Given the introduction of Courant-Snyder invariant, the beam emittance can be defined as the average of the Courant-Snyder of the whole ensemble

$$\bar{\epsilon}_{CS,x} = \frac{1}{N} \sum_{i=1}^N \epsilon_{CS,x,i}$$

so that it will be also invariant. As indicated by Equation (4.2), the computation of emittance cannot be only based on x and x' . The inclusion of the optics function is necessary, as it is in the average definition. However, it is noteworthy that when the beam is matched with the optics function at all instants t , meaning that the normalised particle distribution in the ensemble remains statistically invariant under rotation, the optics function is no longer necessary to define the statistical emittance. In this scenario, the average statistical emittance $\bar{\epsilon}_{CS,x}$ converges to

$$\epsilon_{rms,x} = \sqrt{\det \begin{pmatrix} \bar{x}^2 & \bar{x}x' \\ \bar{x}x' & \bar{x}'^2 \end{pmatrix}} = \sqrt{\bar{x}^2 \bar{x}'^2 - \bar{x}x'^2} \quad (4.3)$$

For more details about this definition of matched beam distribution, it is possible to read ([St20]). It is crucial to recognize that the provided definition of emittance refers to the second central moment, as represented in Figure (4.2).

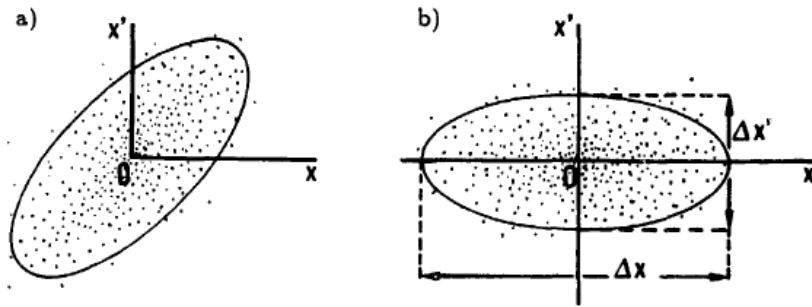


Figure 4.2: Beam particles represented by point in the (x, x') trace space, where on the left there is the tilted ellipse shape while on the right it is upright, ([Bu94]).

It is feasible to prove that the statistical emittance, specifically in instances when the machine lattice is linear, remains constant in accordance with the motion equation. Considering that, LHC can be considered a linear lattice, this assertion of invariance is particularly promising. For further elaboration, interested readers may refer to ([Bu94]).

However, this linearity hypothesis on the motion can be relaxed.

In situations when the number of particles, denoted as N , is quite large in a tiny volume, the state of the particle beam can be described by a phase density function, denoted as

$f(\vec{q}, \vec{p}, t)$. This function allows us to determine the number of particles, dN , present within an infinitesimally small volume at a specific time, t :

$$dN = f(\vec{q}, \vec{p}, t) d^3q d^3p \quad (4.4)$$

There exists a significant theory concerning the temporal evolution of a system in an infinitesimal volume, that holds crucial importance in generalizing the concept of emittance. The evolution of a physical system can be conceptualized as a mapping from the phase space onto itself, establishing a connection between two representative points at distinct moments in time.

Theorem 4.1.1. *Liouville’s theorem:* *When this evolution is governed by Hamiltonian forces, the density function represented in (4.4) is constant when measured at the position of a state moving through phase space*

$$f(\vec{q}(t), \vec{p}(t), t) = f(\vec{q}(t'), \vec{p}(t'), (t'))$$

if the reader wants more details about the connection of this theorem and the beam dynamics can read ([So88]).

From this theorem, the movement of the representative points within a 6-dimensional phase space, as described before, can be compared to the motion of an incompressible fluid. This analogy arises from the similarity between the phase density and the mass density of the fluid.

Specifically, while taking into account the unique form of the equation of motion in (4.1), the Liouville’s theorem applied to our phenomenon turns in stating that the volumes encompassed by a given contour remain unchanged under the mapping function.

In particular, the theorem posits that the “local” beam emittance defined as the volume that contains all the particles in the phase space, considering that these particles are clustered in an infinitesimal volume, remains constant during the beam’s evolution either if the motion is linear or not. This principle applies whatever it is the Hamiltonian system time-invariant.

4.1.2 The importance of emittance

The definition of emittance is important, as indicated by the concise explanation provided below:

- Firstly, the efficiency of collisions is strictly connected with the beam envelope. The objective is to minimize the latter in order to maximize the luminosity at the interaction points, in order to have “better” bunch collisions. This means reducing the beam size to an ideal point-like configuration, thereby enhancing the collision rate. Additionally, this reduction in beam size serves to improve the resolution of fixed target experiments conducted along the ring. As anticipated, achieving this objective is not a simple task, as there are numerous effects that are adverse to our objective of confining $\approx 1e11$ particles within a limited space. As a type of example, the typical value for the transverse emittance is $\approx 2.3 \times 10^{-6}$ m.
- Secondly, it is one integral of motion of the system (4.1) in linear accelerators. This plays a crucial role as it implies that this number, regardless of the circumstances, must remain constant within the accelerator. Recognizing this aspect is essential as it serves as the initial step for further calculations aimed at determining other parameters. Unfortunately, achieving this assertion in practical terms is unfeasible within

real accelerators due to the intricate nature of various effects. In particular, the emittance typically experiences an increase known as “emittance blow up”, which is often attributed to noise sources, which contradicts the initial objective. For further elaboration on this topic, interested readers may refer to the work ([Ku16]) for additional information.

In the context of particle acceleration, it is crucial to replicate a near-zero pressure environment in order to facilitate the attainment of light speed by the particles. To do this, a vacuum chamber has been constructed within the vacuum chamber, containing two distinct levels of vacuum. Understanding the spread of beam dispersion is crucial in determining whether the beam is compatible with the dimensions of the vacuum chamber. In the scenario of “emittance blow up”, it is imperative to comprehend the dispersion of beam spread. This understanding is pivotal in assessing the compatibility of the beam with the dimensions of the vacuum chamber and establishing the temporal feasibility of beam utilisation during the operational phase.

- Thirdly, the detectors possess limitations in their ability to examine all potential particle collisions. In fact, they employ a discerning logic that selectively triggers and filters certain collisions based on certain criteria, so enabling the analysis of only the most significant and relevant collisions. In the context of “emittance blow up,” it is possible for particles to undergo collisions with various facilities within the pipe, distant from the interaction point. As a result, secondary particle showers may reach the detector, leading to the triggering of unproductive collisions. This, in turn, can adversely impact the efficiency of the detector analysis. This phenomenon isn’t going to have a significant impact on the lifetime of the beam during the operational period. However, it may compromise the efficacy of the goal analysis conducted by the accelerator.

4.1.3 Emittance measurement

Accurately measuring the longitudinal emittance can be achieved, relatively easily, by employing a monitoring technique to attain the desired value ([Ca09]). Conversely, measuring the transverse beam emittance poses significant challenges due to the temporal variation of the signal along the longitudinal axis in comparison to the stationary nature of the transverse axis during the run. Additionally, the concentration of beam energy in the transverse plane (approximately in the order of GeV) is significantly smaller than that along the longitudinal axis ≈ 7 TeV.

To determine the transverse beam emittance, it is necessary to acquire the transverse beam profile. This can be accomplished by accelerating the beam towards a stationary target to gather the desired data. However, this approach results in the loss of the beam, thereby compromising the primary objective of obtaining collisions and luminosity.

Unfortunately, performing an accurate measurement of the transverse emittance is challenging due to intrinsic limits of the adopted techniques, of the working assumptions, and to the potential alterations to the beam’s transverse emittance caused by certain methodologies.

The development and improvement of the transverse emittance measurement techniques has been the subject of extensive research for several decades, dating back to the 1950s. A substantial body of literature exists, encompassing various methodologies and techniques.

Currently, the Wire Scanners are a typical instrument used for emittance measurement in the accelerators, such as the SPS at CERN. The transverse beam profile and beam emittance are determined by analysing the shower of secondary particles resulting from the interaction between a rapidly moving thin wire and the beam. This interaction occurs within a vacuum chamber, and the detection of the secondary particles takes place outside the chamber on an

assembly. From this measurement it is possible to reconstruct the transverse beam profiles, and so the transverse beam emittance. This particular approach is widely employed and known for its high level of precision. However, a drawback of this strategy is that even when the wire used is extremely thin, it still has an impact on the beam structure. This impact, while minor, results in a modification of the beam profile, so altering the actual emittance value from that moment onward.

An alternative methodology, commonly employed in the LHC, involves the utilization of the Synchrotron Light monitor. This monitoring technique uses the Synchrotron radiation emitted by the particle beam. The emitted radiation of the beam is then directed towards a specialized stationary detector through the use of lenses. This detector, by examining the properties of the light, such as its intensity and wavelength, can determine the transverse emittance of the beam. For further elaboration on the aforementioned instruments, interested readers may refer to the work ([Ku16]).

However, a definitive method for achieving transverse emittance that is as pristine as the longitudinal counterpart has not yet been established. The most prevalent approach employed in the accelerators thus far has been outlined. However, an extensive body of literature exists on these analyses, encompassing a wide range of methodologies. Only a select few have been mentioned here:

- Quadrupole Scans ([Mi03])
- Three-gradient measurements ([Ol13])
- IPM(Ionization Profile Monitor) ([St15])

The objective of this part of the project was to explore an alternative methodology for determining the transverse emittance, without relying on any specific tool employed in an accelerator, but rather employing a numerical approach starting from the measured luminosity. The objective is to invert the luminosity integral, so the transverse emittance of the beam is obtained from the luminosity, which represents one of the most accurate measurement in an accelerator. The relationship between luminosity and transverse beam emittance is clearly apparent from a technical point of view, while its mathematical representation will be derived in the following section.

4.2 Luminosity model computation

The purpose of this section is to define a general formula for the computation of luminosity which contains some intrinsic accelerator aspects (crossing angle, offset, hourglass effect...) but, with the aid of several assumptions, can address and compute luminosity in an optimized manner. It is imperative for the purpose of this section of the thesis that this model satisfies these two prerequisites, and the reasons for this will be elucidated in Section XYZ.

4.2.1 The luminosity integral

Luminosity is derived as a scalar output resulting from a convolution of multiple factors that represent several features of the beam during collision. This notion encompasses various aspects, including beam emittance, energy, frequency of beam oscillation, number of bunches, and the number of collisions between individual bunches, etc There exist two discernible categories of luminosity, specifically the luminosity linked to a beam interacting with a fixed target and the luminosity related with the collision of two beams. The aim of this thesis is to conduct a comprehensive study of the second case.

In this particular scenario, both beams operate simultaneously as the “target” and “incoming” beam, as represented in Figure (4.3).

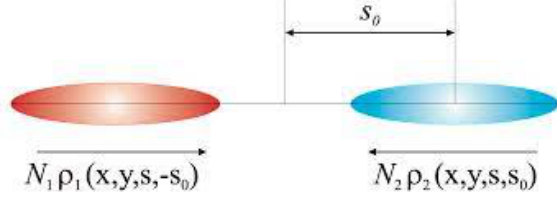


Figure 4.3: The bunches before collision and thus displaced, in which the target of one beam is precisely the position of the other. It is evident that, considering that in the accelerator all the bunches are in a row, there will be other collisions away from the interaction point, ([Co14]).

The analysis of beam density distribution along the three axes plays a critical role in understanding the dynamics of this specific collision scenario, especially when compared to the scenario involving a fixed target. Moreover, it is important to acknowledge that the two beams under consideration are not in a stationary state, but rather in motion with respect to each other with a velocity comparable to that of light. Consequently, the degree of their intersection is impacted not solely by the longitudinal positioning of the bunches, but also by the time duration in which they intersect. The determination of luminosity in this context requires the calculation of an overlap integral across the three spatial dimensions and the temporal dimension. Consequently, it is reasonable to assert that

$$L = N_1 N_2 n_b f_r \sqrt{(\vec{v}_1 - \vec{v}_2)^2 - \frac{\vec{v}_1 \times \vec{v}_2^2}{c^2}} \iiint \int_{-\infty}^{\infty} \rho^{B1} \rho^{B2} dx dy dz dt \quad (4.5)$$

where ρ denotes the normalized beam distribution in a four-dimensional space. Let N represent the quantity of particles per bunch, $\vec{v}$ denote the velocity, and n_b indicate the number of bunch collisions in the interaction point (IP). The variable f_r represents the frequency of collisions, whereas c denotes the amplitude of light speed. Lastly, the square element is referred to as the Moeller factor.

This is the most commonly employed and straightforward approach to delineate the luminosity in the context of colliding beams within a circular or linear accelerator. As the inclusion of hypotheses is introduced into this definition, the mathematical formulation progressively gets more intricate, especially when applying this definition to the specific context of the LHC for determining the track of the particle beam within the accelerator ring

In order to specify the formula and render it analysable, let's introduce some hypothesis. However, before their assumption, it is necessary to establish some parameters that will play a crucial role in the ultimate depiction of the model, focusing mostly on the LHC structure.

4.2.2 Introducing the parameters

Crossing angles

The utilization of beam velocity data is evident in the computation of luminosity. Specifically, the exact velocities of the two beams in the LHC are not known, but they can be determined from the energy of the beams in (eV), as set by the accelerator, by calculating the Lorentz factor derived from the Lorentz transformations of special relativity.

One essential attribute of the LHC, which is pivotal for determining the track of the particle beams within the accelerator ring and so also for their velocity, are the crossing angles. The LHC accommodates nearly 3000 tightly packed bunches. If the two beams were to remain confined within the same pipe, each bunch from one beam would collide with all the bunches from the other beam at various locations along the ring where no dedicated detector is present to capture the secondary particles, as shown in Figure (4.3). Consequently, these particle collisions within the pipe give rise to multiple mechanical issues within the accelerator. Hence, in order to mitigate the occurrence of multiple undesired collisions, the two beams in the LHC are directed through separate pipes. As the beams approach the interaction point, magnets are employed to deflect them at a crossing angle of approximately $\approx 160 \mu \text{ rad}$, as shown in Figure (4.4). This strategic deflection ensures that the beams collide exclusively at the designated location of the detector, effectively preventing any other incidental collisions.

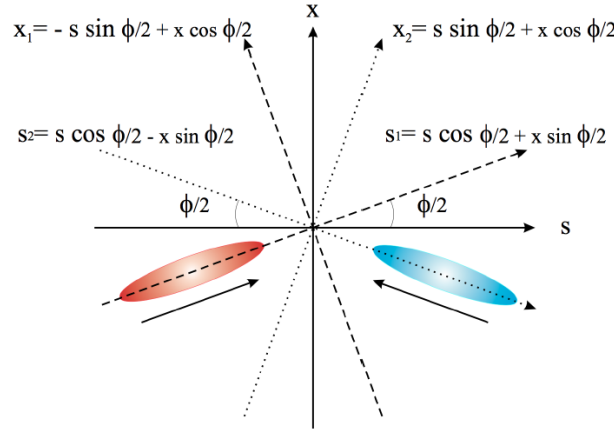


Figure 4.4: The bunches before collision and thus displaced, rotated by a crossing angle, in such a way that the bunches don't collide Head-On apart from the interaction point. There will be electromagnetic forces between the bunches apart from the interaction point, called Long-Range collision, that are particularly smaller than the Head-On ones, ([Co14]).

The crossing angle in the Figure (4.4) was implemented along the x-axis. This approach can be readily expanded, particularly in the context of the LHC, where the crossing angle usually vary from zero in one plane for each detector. The evaluation of the overlap integral occurs in both the x and y dimensions. However, altering the collision angle results in a transition of the reference system to a rotated system. For further elaboration on the computation, interested readers may refer to the works of ([Mu06]) and ([Co14]) for additional information. Drawing inspiration from various literary references and looking into the functioning of the LHC, it is plausible to highlight, assuming one hypothesis, the inherent symmetry of the problem. This assumption, which does not yield any discernible disparities, states that the overall crossing angle Φ is composed of two rotations, namely $\Phi/2$ and $-\Phi/2$, for each beam within the desired plane, as shown in Figure (4.4).

Offset

A modification of the previous structure is required in instances where the beams do not collide directly, but instead have a slight transverse offset, as shown in Figure (4.5). This scenario can be analysed with or without a crossing angle. The objective of this arrangement is to induce a distinct collision pattern between the bunches, wherein the core of one bunch's

particle density collides with the tail of the other bunch.

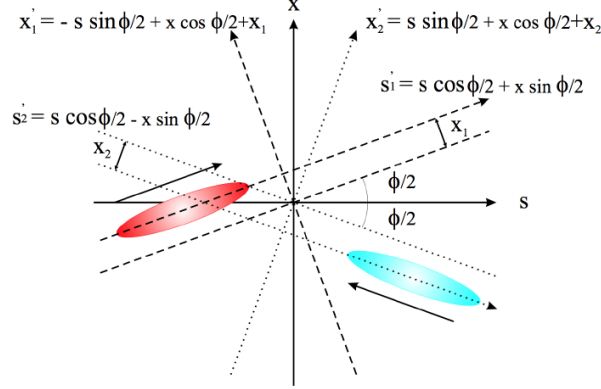


Figure 4.5: The bunches before collision and thus displaced, rotated by a crossing angle and shifted in the transverse plane, in this situation the target of the bunches of one beam isn't precisely the other bunch, and the collision Head-On didn't happen any more, in particular there will be a collision between a tail of one bunch with the core of the other one, ([Co14]).

For further elaboration on the computations in this particular example, interested readers may refer to the works of ([Mu06]) and ([Co14]) for additional in-depth information. As previously examined, and in order to highlight the structural symmetry, it can be postulated in this instance that the overall offset (X_1 for B1 and X_2 for B2) consists of two opposite offsets, namely $\mu_0/2$ and $-\mu_0/2$, for each beam within the designated plane.

Twiss parameters

From a theoretical point of view, it is widely acknowledged that the beam emittance holds significant importance as a fundamental invariant within the system. It should be noted that the variances ($\sigma_x, \sigma_{x'}$), that represent the beam size and the beam divergence, are not fixed quantities. Specifically, they are subject to variation along the beam line, as they are influenced by the magnet optics. To mitigate the influence of the magnets on the beam characteristics, it is feasible to perform a normalization procedure:

$$\sigma_x^2(s) = \epsilon_x \cdot \beta_x(s)$$

$$\sigma_x'^2(s) = \epsilon_x \cdot \gamma_x(s)$$

$$r\sigma_x^2(s)\sigma_x'^2(s) = -\epsilon_x \cdot \alpha_x(s)$$

where r is the correlation coefficient. These parameters that encompass all the influence of the magnets are the same as the ones defined in the (4.1.1)

Hourglass effect

As previously mentioned, the variances in beam density within the specified plane can exhibit shifts. We have accurately described the influence of the magnets on this phenomenon to emphasize the theoretical invariance of the transverse emittance. In particular, the β function exhibits variation in proximity of the interaction point:

$$\beta(s) = \beta^* - 2\alpha^*z + \gamma^*\beta^* \quad (4.6)$$

where the notation involving the star signifies the evaluation of the function at a curvilinear abscissa of zero (i.e. IP). This suggests a correlation between the densities of transverse and longitudinal beams. The phenomenon described is commonly referred to as the Hourglass effect, which is attributed to the parabolic curvature of the β function in proximity of the interaction point, as represented in the Figure (4.6).

The actual impact of this phenomenon would result in a modification of the luminosity factor when the bunch length is comparable to the β^* . Otherwise, it can be regarded as insignificant. Specifically, within the context of the LHC, the length of the particle bunch is denoted as $\sigma_z = 7.55$ cm, whereas the beta function is represented as $\beta^* = 30$ cm. Consequently, it is crucial to account for this particular phenomenon. For further elaboration on the computations, interested readers may refer to the works of ([Mu06]) and ([Co14]) for more information.

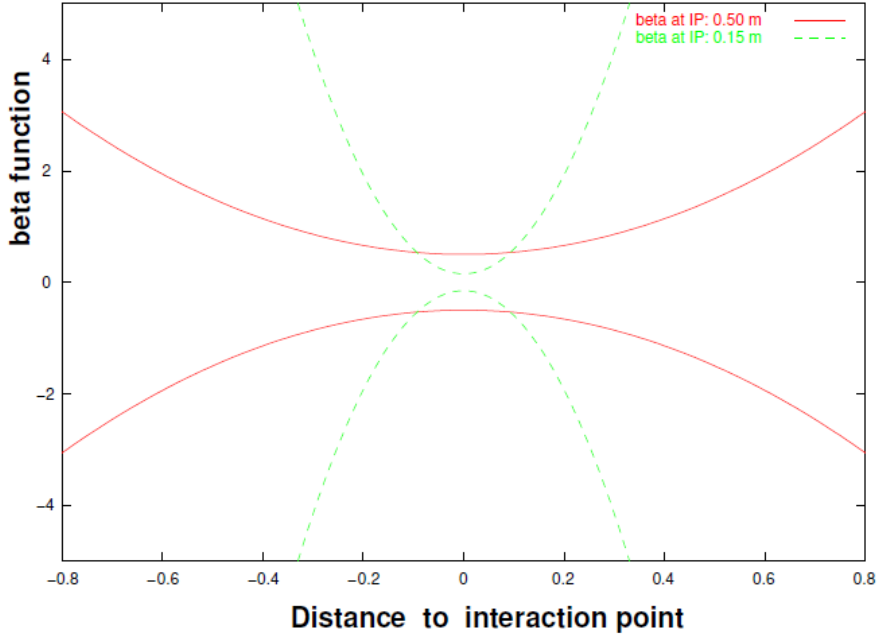


Figure 4.6: Illustration of the hourglass effect showing how varies the beta function, considering alpha null, in the longitudinal axis around the IP for two different values of β^* with the LHC nominal bunch length of 7.55 cm, ([Co14]).

The many characteristics and effects that must be taken into account for a comprehensive description of the problem have been described. By incorporating all the necessary assumptions, it is feasible to establish a model for the luminosity of two colliding beams.

4.2.3 Model and hypothesis

The computation involved in the luminosity formula presents challenges, apart from the convolution and quadruple integral. The difficulties arise from the complex nature of the intrinsic properties of the beam density. These complexities make it impossible to obtain a closed-form solution, and in some cases, it may be neither feasible to reduce the dimensions of the integral.

Thus, in order to enhance the comprehensibility and accessibility of the formula, it is necessary to establish certain assumptions beforehand.

Factorization of distributions

Firstly, considering the structure of equation (4.5), it is more convenient to assume the hypothesis of distribution factorization. This hypothesis suggests that the distributions along the various geographical axes are independent and not correlated.

$$\rho^B = \prod_{i \in \{x, y, z\}} \rho_i^B, \quad \text{with } B \in \{B1, B2\}$$

from this hypothesis, it is no longer feasible to take into account any sort of coupling effects in the transverse plane ([Br91]). Consequently, the formula (4.5) becomes

$$L = N_1 N_2 n_b f_r \sqrt{(\vec{v}_1 - \vec{v}_2)^2 - \frac{\vec{v}_1 \times \vec{v}_2^2}{c^2}} \iiint_{-\infty}^{\infty} \rho_x^{B1} \rho_x^{B2} \rho_y^{B1} \rho_y^{B2} \rho_z^{B1} \rho_z^{B2} dx dy dz dt \quad (4.7)$$

Transverse dependence with respect to z

Furthermore, based on the paraxial approximation, which assumes that particle trajectories are in proximity to the optical axis (i.e. curvilinear abscissa along the ring) and that deviations from the axis are negligible compared to the ideal trajectory, the equation of motion can be simplified. This simplification allows for the expression of these deviations as linear transformations ([Hi21]). From this, without considering any other possible mechanism in the accelerator, it's feasible to assume that

$$\begin{aligned} \rho_x^B &= \rho_x^B(x, z), \quad \text{with } B \in \{B1, B2\} \\ \rho_y^B &= \rho_y^B(y, z), \quad \text{with } B \in \{B1, B2\} \\ \rho_z^B &= \rho_z^B(z, t), \quad \text{with } B \in \{B1, B2\} \end{aligned}$$

For additional clarification on this idea, interested readers might refer to ([Sy07]). Based on the aforementioned hypothesis, since ρ_z is independent of both x and y, the luminosity formula (4.7) can be written like this:

$$L = N_1 N_2 n_b f_r \sqrt{(\vec{v}_1 - \vec{v}_2)^2 - \frac{\vec{v}_1 \times \vec{v}_2^2}{c^2}} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} \rho_z^{B1} \rho_z^{B2} dt \int_{-\infty}^{\infty} \rho_x^{B1} \rho_x^{B2} dx \int_{-\infty}^{\infty} \rho_y^{B1} \rho_y^{B2} dy dz \quad (4.8)$$

Beam density as normal distributions

Thirdly, it is worth noting that the beam distribution in each axis can be reasonably modelled as a normal distribution. However, it is important to acknowledge that this approximation is not entirely accurate due to the dynamic nature of the beam shape during the operation of the LHC. Recent research ([Pa20]) has reported that the beam distribution becomes more intricate, following a Q-Gaussian distribution. Nevertheless, for the purposes of this project, the analysis is confined to the assumption of a normal distribution:

$$\rho_i^B = \frac{1}{\sqrt{2\pi}\sigma_{i,B}} e^{-\frac{(i-\mu_{i,B})^2}{2\sigma_{i,B}^2}}, \quad \text{with } B \in \{B1, B2\} \text{ and } i \in \{x, y, z\}$$

where the $\mu_{z,B} = \beta_r ct$, with β_r that represents the relativistic beta from the Lorentz transformations.

The primary objective of this assumption is to simplify the 4-dimensional integral in x , y , z , and t into a 1-dimensional integral in z . This is achieved by analytically solving the integral in time and in the transverse plane, hence improving the efficiency of the numerical integration for the unsolvable portion.

After performing the process of analytical integration, it is necessary to describe the mean and variance of the beam distribution in the various planes as a function of z . Subsequently, the final integral needs to be numerically integrated.

The closed form of the transverse integral on the x -axis can be represented based on the results obtained through analytical computations. It is worth noting that the y -axis has a similar structure.

$$x_{factor} = \frac{\sqrt{2}e^{-\frac{\frac{\mu_{x1}^2}{2} + \mu_{x1}\mu_{x2} - \frac{\mu_{x2}^2}{2}}{\sigma_{x1}^2 + \sigma_{x2}^2}}}{2\sqrt{\pi(\sigma_{x1}^2 + \sigma_{x2}^2)}}$$

as this, it's it is feasible to depict the outcome factor of the t -integration

$$t_{factor} = \frac{\sqrt{2}e^{-\frac{z^2\left(-\frac{\beta_{r1}^2}{2} + \beta_{r1}\beta_{r2} - \frac{\beta_{r2}^2}{2}\right)}{\beta_{r1}^2\sigma_{z1}^2 - \beta_{r2}^2\sigma_{z2}^2}}}{2c\sqrt{\pi(\beta_{r1}^2\sigma_{z1}^2 + \beta_{r2}^2\sigma_{z2}^2)}}$$

Finally, the formula for the computation of the luminosity is

$$L = N_1 N_2 n_b f_r \underbrace{\frac{\sqrt{c^2(\vec{v}_1 - \vec{v}_2)^2 - \vec{v}_1 \times \vec{v}_2}}{2c^2}}_{\text{Moeller efficiency}} \int_{-\infty}^{\infty} e^{-\frac{\frac{\mu_{x1}^2}{2} + \mu_{x1}\mu_{x2} - \frac{\mu_{x2}^2}{2}}{\sigma_{x1}^2 + \sigma_{x2}^2}} e^{-\frac{\frac{\mu_{y1}^2}{2} + \mu_{y1}\mu_{y2} - \frac{\mu_{y2}^2}{2}}{\sigma_{y1}^2 + \sigma_{y2}^2}} e^{-\frac{z^2\left(-\frac{\beta_{r1}^2}{2} - \beta_{r1}\beta_{r2} - \frac{\beta_{r2}^2}{2}\right)}{\beta_{r1}^2\sigma_{z1}^2 + \beta_{r2}^2\sigma_{z2}^2}} \frac{1}{\sqrt{2\pi^{\frac{3}{2}}\sqrt{(\sigma_{x1}^2 + \sigma_{x2}^2)}\sqrt{(\sigma_{y1}^2 + \sigma_{y2}^2)}\sqrt{(\beta_{r1}^2\sigma_{z1}^2 + \beta_{r2}^2\sigma_{z2}^2)}}} dz \quad (4.9)$$

Moeller factor as constant

Finally, it can be postulated that the velocities of the two beams ($\vec{v}_1, \vec{v}_2$) are not influenced by changes in spatial or temporal variables. In the context of ultra-relativistic conditions, where the velocities of the two beams satisfy $\vec{v}_1 = -\vec{v}_2$ and $|\vec{v}_1| = |\vec{v}_2| = c$, it has been seen that the Moeller efficiency may be well approximated as 1. Consequently, the final formula for the calculation of luminosity can be expressed as follows

$$L = N_1 N_2 n_b f_r \int_{-\infty}^{\infty} e^{-\frac{\frac{\mu_{x1}^2}{2} + \mu_{x1}\mu_{x2} - \frac{\mu_{x2}^2}{2}}{\sigma_{x1}^2 + \sigma_{x2}^2}} e^{-\frac{\frac{\mu_{y1}^2}{2} + \mu_{y1}\mu_{y2} - \frac{\mu_{y2}^2}{2}}{\sigma_{y1}^2 + \sigma_{y2}^2}} e^{-\frac{z^2\left(-\frac{\beta_{r1}^2}{2} - \beta_{r1}\beta_{r2} - \frac{\beta_{r2}^2}{2}\right)}{\beta_{r1}^2\sigma_{z1}^2 + \beta_{r2}^2\sigma_{z2}^2}} \frac{1}{\sqrt{2\pi^{\frac{3}{2}}\sqrt{(\sigma_{x1}^2 + \sigma_{x2}^2)}\sqrt{(\sigma_{y1}^2 + \sigma_{y2}^2)}\sqrt{(\beta_{r1}^2\sigma_{z1}^2 + \beta_{r2}^2\sigma_{z2}^2)}}} dz \quad (4.10)$$

The dependency of the emittance in this final luminosity formula is not immediately evident. Consequently, the task of inverting the formula to get the transverse emittance appears to be unachievable. To determine the relationship with the emittance, it is necessary to establish the functional expressions of the average and standard deviation of the transverse beam distributions with respect to z .

Specifically, the mean value exhibits a z-linear trajectory as a result of the influence of the crossing angle

$$\mu_x(z) = \mu_{x0} + p_x(z=0)z \quad (4.11)$$

where μ_{x0} represents the average value of μ_x at the interaction point, denoted as $z = 0$. This average is associated with the offset of the beams.

Meanwhile, the standard deviation of a given quantity is influenced by two distinct factors, namely the betatronic component denoted as σ_β and the off-momentum contribution denoted as $\sigma_{\frac{\Delta p}{p}}$. The first contribution is influenced by the magnetic quadrupole lenses' strength, which selectively focus or defocus particles depending on their transverse position. In contrast, the second factor describes the variations in particles' transverse position resulting from differential magnet-induced kicks due to their slightly disparate energy levels. For further elaboration, interested readers may refer to ([Hi21]) for additional information.

The final representation of the standard deviation is

$$\sigma_x(z) = \sqrt{\sigma_{x,\beta}^2(z) + \sigma_{x,\frac{\Delta p}{p}}^2(z)}$$

with

$$\sigma_{x,\beta}(z) = \sqrt{\frac{\beta(z)\epsilon_{x,n}}{\beta_r\gamma_r}}$$

where it's important to differentiate between the relativistic Lorentz parameter, denoted as β_r , and the β function, which is associated with accelerator optics, and

$$\sigma_{x,\frac{\Delta p}{p}}(z) = dx(z) \frac{\Delta p}{p}$$

This depiction provides a clearer understanding of the desired dependence and the underlying structure of the transverse emittance with respect to the standard deviation.

4.2.4 Cost computation optimization

The calculation of luminosity, based on the model, involves performing a numerical integration on the final integral with respect to the longitudinal variable. This is necessary as a closed-form solution does not exist. The resulting implementation in Python was discovered ([PythonLumi]). This Python implementation utilises the widely-used tool “`scipy.integrate.quad`” to calculate the numerical integral. The “`scipy.integrate.quad`” function employs quadrature formulas to approximate the desired values. To obtain this estimation, the tool evaluates the integrand multiple times, typically around $\sim 1e2$ evaluations. It is important to note that these evaluations are not instantaneous.

This implies that employing this model to depict luminosity in an iterative manner or in the context of data analysis, or representing this model as a function of certain parameters, may not be instantaneous, mostly due to the intricate nature of computing the integrand function.

To enhance the efficiency of this computation, the decision was made to utilise Numba. Numba is a Just-In-Time (JIT) compiler for the Python programming language, specifically developed to enhance the performance of numerical computations, particularly those involving NumPy arrays, functions, and loops. Numba accomplishes this objective by transforming Python code into machine code that is highly optimised, hence enhancing performance without necessitating code rewriting. Upon the initial execution of the Python compiler, it

compiles the provided code and stores it in the memory. Consequently, subsequent executions of the same code, even with little modifications, will not necessitate recompilation of all functions, as the compiler retains the compiled version in memory. If the reader desires further information, they may refer to the source cited as ([Numba]).

The first code implemented in ([PythonLumi]) was modified to achieve the same result while employing a different computational strategy. Specifically, the code was re-adapted to optimise its performance using Numba. As a consequence, the computation was much enhanced. Without the implementation of Numba, the integration process took $\approx 1 \times 10^{-2}$ s. However, with the inclusion of Numba, the initial compilation time increased to around one second, but the subsequent integration time decreased significantly to $\approx 1 \times 10^{-4}$ s. This optimization in computational cost is evident when performing numerous integration, as shown in the Figure (4.7).

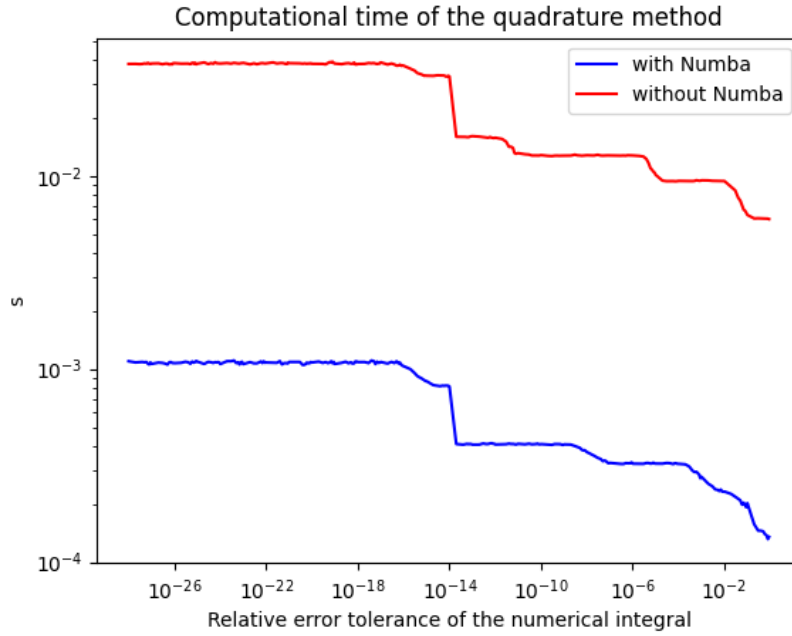


Figure 4.7: Representation of the different average, over 2000 cases, for the computational time of the quadrature method implemented in SciPy for the numerical integral in the luminosity formula between the case in which the integrand is defined with and without Numba. The plots are represented over the relative error tolerance of the numerical integration.

In the figure, it has been performed the average of time over 2000 integration, restarting the python kernel after transitioning from one value of the relative error tolerance of the quadrature method to the next. By restarting the kernel, it becomes feasible to observe the impact of the initial compilation. This is evident as there exists a discernible difference, which is less than two orders of magnitude, between the two plots.

Chapter 5

Inversion problem: emittance from luminosity

The objective of this chapter is to employ the mathematical model proposed in the preceding chapter to perform an inversion and derive the transverse emittance from the luminosity. This goal is particularly interesting, as we aim to derive specific variables from a scalar that is generated by the convolution of multiple factors and parameters. Firstly, it is apparent that the problem at hand exhibits ill conditioning. Indeed, it is already evident, in a simplification, that it is not feasible to reconstruct the original three addends solely from the sum of three numbers. In this case, the final scalar represents the result of a notably more intricate computation. However, the ultimate objective remains unchanged: to address and resolve a problem that is characterised by ill conditioning.

5.1 Non-linear system to obtain the emittance

The previous Chapter (4.1) provided a comprehensive definition of emittance, highlighting its meaning and elucidating its conceptual representation. Additionally, the chapter also presented a review of relevant literature pertaining to the measurement of emittance. The measurement of emittance poses a significant challenge in the field of accelerator physics, particularly in the context of transverse emittance. Moreover, many of the existing measurement techniques employed in accelerator physics have the potential to perturb the beam, thereby altering the emittance value being measured. Furthermore, it is worth noting that the majority of commonly employed measuring techniques exhibit a relative estimation error on the order of 10%. Consequently, the accurate determination of emittance remains an unresolved issue in this field.

Despite the inherent complexity and non-linearity of the model, the objective is to effectively inverse it by leveraging the intricate structure of the machine under analysis. This involves attempting to appropriately and realistically use the flexibility of the detectors, even if just in simulated scenarios.

Upon concluding this introductory section, readers may find themselves contemplating the methodology behind the use of luminosity and wondering the underlying motivations for using it with respect to other factors on which the emittance has a significant impact.

5.1.1 The reason why it has been used the luminosity

The luminosity holds significant importance in accelerators, as it is as a scalar quantity that reflects the efficiency of the accelerator. It quantifies the amount of information that

can be obtained from the collisions of bunches throughout the entire operational period. By convoluting diverse factors within a single value, the luminosity compactly represents the intricate nature of collision phenomena.

Properly for this reason is one of the most accurate measurements in the LHC, it is an indirect measurement derived from the secondary particles produced during the collision of particle bunches. Various detectors are employed for this purpose, with numerous researchers dedicated to studying the most precise methods for its measurement. Extensive literature exists on this intricate measurement process. Hence, the measurement in question holds significant precision within the LHC considering its magnitude, estimated to be approximately $e34$, indeed this measurement is regarded accurate up to the third significant digit (denoted as $1e - 3$).

Numerous parameters influence this value, with some being precisely measured, others manually controlled by the machine, and finally, the emittances, which are known to have the lower accuracy measurements.

Hence, our objective is to invert a challenging non-linear model, which arises from the convolution of a quadruple integral and that is influenced by numerous parameters. Furthermore, the objective is to acquire a four-dimensional vector, that represents emittance values for each plane and each beam, from a single scalar. This is a priori, even in the most simple computation possible (a sum of different values), an ill conditioned problem

So, finally the primary idea is to identify a method for strategically conditioning the problem, with the aim of becoming the impossible (an ill conditioned problem with multiple solutions for the inversion problem) difficult (a well conditioned problem particularly complex with inversion well-defined and unique). The primary analysis revolves around the appropriate manner in order to condition the problem.

5.1.2 The idea behind this approach

It is not feasible to derive a vector from a single scalar, namely the luminosity, by inverting its model. The fundamental idea of this approach is to examine the collision from various perspectives and viewpoints shifting some parameters from their initial values, in order to establish a system which includes distinct and independent equations, which collectively enable the determination of the desired unknowns. The equations presented herein will be formulated from the difference between the luminosity value and the corresponding model with that specific choice of parameters, named configuration system. In order to have independent equations, it has been used the nonlinear dependence of luminosity on various parameters, and by shifting these parameters using the flexibility of the detectors, this collision has been analysed from multiple, non-redundant perspectives.

Inspiration of this approach

The motivation for this analysis derives from the vast amount of literature exploring the use of luminosity values in various configurations for beam diagnostics, as in ([Hos18]). The primary difficulty in this study lies in rejecting a main hypothesis, which is commonly seen in the existing literature, that assumes an equal value of emittance across various beams and/or planes.

The methodology employed in this study can be classified as a form of perturbation theory. Its objective is to establish a more precise characterization of the relationship between luminosity and four distinct values of emittance. This is achieved by manipulating certain parameters, creating a system of equations, to highlight the underlying numerical dependencies of the model. By employing the Nonlinear Least Squares method of SciPy ([LS])

(denoted LS) on the system, the model's various dependencies can be discerned, allowing for the determination of the emittances.

Some challenges may arise in certain cases from the inherent complexity of the problem, which exhibits numerous symmetries. Consequently, it becomes difficult to manipulate the parameters in a strategic manner to effectively condition the problem. This difficulty is particularly pronounced in the most general scenario, where the objective is to obtain four distinct emittances. In this case, representing the penalty function of the LS algorithm as a function of four variables is particularly intricate.

The concept of manipulating certain parameters to generate diverse equations for the same collision is derived from a method studied by a Nobel Prize that is the Van der Meer scan method, which was developed in 1968 and continues to be employed in the LHC for luminosity calibration purposes. This method involves transversely sweeping the beams across each other ([Ba21], [Wh10]).

Choice of the parameters

There exist various parameters that do not exhibit a linear relationship with luminosity, including the offset, crossing angle, longitudinal standard deviation, β function, and α function. Specifically, considering the dependence of the parameters in the model, it is not necessary to have the maximum degree of freedom when adjusting them. For instance, if there is a need to modify the offset of the x-axis for the B1 component of the quantity C without altering the offset of the x-axis for B2, the same outcome can be achieved by shifting the B1 component of C/2 and the B2 component of -C/2. Due to the presence of certain shifts that result in redundant information and the potential interconnection between the shifting parameters of the beams, a decision has been made to consistently change both beams simultaneously. Specifically, altering the offset or crossing angle of B1 along a given axis results in a corresponding shift in the opposite direction on the same plane for B2. Conversely, when it pertains to the longitudinal beam size and the Twiss parameters, modifying the parameter for B1 by C/2 produces an equivalent shift of C/2 for B2.

Furthermore, among the several potential parameters that can be shifted, certain ones are favoured due to the simplicity of modification within the machine and their greater compatibility for precise adjustment. The offset and crossing angle are relatively easier to modify due to their precise calibration by the machine. Subsequently, the longitudinal beam size is considered, followed by the Twiss parameters.

To enhance our comprehension of the underlying physics and the luminosity model, we conducted experiments on various parameter combinations. Specifically, we will primarily present the results pertaining to shifting of offset and crossing angle, while also considering other cases.

To gain a distinct perspective on the collision during the shift on the parameter, and taking into account that luminosity is precisely measured up to the 3rd significant digit ($1e-3$ of its nominal value), an intentional adjustment is made to a specific parameter in both beams. This adjustment is designed to collectively impact the luminosity value by a factor of $1e-2$, thereby enabling a more significant observation of a different measurement value. This substantial alteration in the nominal value provides valuable information within the mathematical framework that necessitates inversion. In the configuration system of the LHC, this adjustment on the parameters that leads to a 1% change in nominal luminosity has been demonstrated to be realisable from a practical and engineering standpoint.

Non-linear Least Square method

To further clarify the subsequent section's results plot, a brief introduction is provided on the non-linear least squares method implemented in SciPy. Our use of this method is intentionally simplified to ensure its intuitive nature.

By using the default parameters (penalty function, method for the Jacobian, method for the minimization) and providing the system of equations that we aim to minimize from an initial guess of the solution, a penalty function is formed:

$$F_{\text{penalty}}(\vec{\epsilon}) = \sum_{i=0}^n f_i(\vec{\epsilon})$$

The penalty function is evaluated at the initial guess, as well as two other points that are slightly displaced along the axis relative to the initial guess. The method evaluates the plane that intersects these points to approximate the “slope” of the penalty function and move along the direction of sharpest fall. The algorithm progresses in that particular direction until it locates the optimal point, characterised by a reduced penalty function. It then iteratively repeats this process, starting from the most recently discovered point, until it finally arrives at the same point as the initial one of that step. For further information regarding the features, parameters, and procedures, interested readers may refer to the source ([LS]).

Based on the concise description provided, it is evident that there is a necessity to employ the numba library in representing the model discussed in section (4.2.4). This is due to the presence of multiple equations and many steps involved in the LS process, resulting in a significant number of evaluations of the integrand. Consequently, this could potentially pose challenges in accurately computing the inverse problem without using Numba, leading to a substantial increase in the time required to resolve it, lasting hours or even days.

5.2 1st inversion problem, without errors

The derived formula for the luminosity exhibits a high level of complexity, incorporating several inherent characteristics of the problem as well as factors related to the flexibility of the machine. In virtue of this, it is imperative to partition the many input variables of the luminosity expression, so achieving a more formal and beneficial framework for the mathematical delineation of the problem.

To do this, the various parameters have been divided into three distinct clusters:

- The unknown variables, referred to as emittances, provide information about the sigma value of the Gaussian beam profile along the axes, ($\vec{\epsilon} = (\epsilon_{x1}, \epsilon_{x2}, \epsilon_{y1}, \epsilon_{y2})$)
- The variable parameters refer to those parameters that may be adjusted on the machine. These parameters are used to achieve different luminosity values for various machine configurations, in order to have a more comprehensive understanding of the collision process, ($p_{\text{var}} = (\mu_{\{x,y,1,2\}}, p_{\{x,y,1,2\}}, \beta_{\{x,y,1,2\}}^*, \alpha_{\{x,y,1,2\}}^*, \sigma_{z\{1,2\}})$), where the notation of the subscript signifies ($v_{\{x,y,1,2\}} = v_{x1}, v_{x2}, v_{y1}, v_{y2}$)
- The constant parameters refer to those parameters that must be set at fixed values in order to ensure the proper functioning of the LHC according to its programmed specifications. These parameters are either non-shiftable or, if they can be shifted, their variations exhibit a linear relationship with the luminosity. Consequently, any shifts in these parameters would not provide meaningful insights into the final system. ($p_{\text{const}} = (n_b, f_r, N_{\{1,2\}}, H_{\{1,2\}})$), where H is the total energy of the beam

$$\mathcal{L} = \mathcal{L}(\vec{\epsilon}|p_{var}, p_{cost})$$

Starting from the initial choice of the parameters (initial configuration system), that correspond to one equation, the primary objective is to identify a shift in diverse components of p_{var} that will result in a modification of the initial luminosity value by 1%. Subsequently, the model will be recalculated using these adjusted values of p_{var} , leading to a distinct equation. Therefore, proceeding in this manner, a system of equations will be generated, each, except the first one related to the initial configuration system, with the same luminosity value but computed using different shifted parameters. The ultimate objective is to employ a Non-linear Least Square method to reduce the complexity of this system, with the aim of determining the desired unknown emittances.

It is noteworthy that the initial values assigned to p_{var} are primarily derived from the actual values employed in the LHC. Consequently, there is no justification for seeking different machine conditioning between the two beams. In this analysis, all parameters will therefore be equally set for both beams, or in some cases, set as opposites. As a result, any shifts observed in the relevant parameters will adhere to this symmetry.

This approach is similar to a simulation, wherein the equations used in the LS method are generated by the identical model. The manner in which this is accomplished is as follows:

- From the selected initial configuration system for a detector, a decision has been made on which parameter to modify, resulting in a new value for that parameter.
- Selecting the four emittances randomly, which are assumed to be unknown. These emittances are then provided as input to the model with the objective of calculating the luminosity.
- The luminosity is acquired and afterwards used in the equation by subtracting it from the model with the selected configuration system.
- The equation is incorporated into the system, subsequently it has been decided another parameter for shifting, and the process restarts from the original stage.

In practise, when implementing this methodology within the accelerator, the control room of the LHC may set the different values for the parameters. Moreover, it is important to note that the luminosity value is derived from detector measurements rather than relying on the model employed for equation composition.

The passage from simulation to a realistic method highlights the presence of some hypotheses that have thus far been omitted, but are nevertheless significant to emphasise:

- The accurate determination of luminosity without of random error: it has been assumed that the value obtained from the model is perfect and devoid of any errors. However, it is important to note that in practise, the luminosity value is obtained through machine measurements, which inherently possess a random error of up to $1e - 3$.
- The accurate determination of parameter values: it is contingent upon the assumption that the values set from the machine control room correspond precisely to those present in the physical accelerator facility. However, it is important to acknowledge that these values may be subject to both random errors arising from the limitations of the measurement tool's precision, and systematic errors resulting from incorrect calibration of the facility.
- The exact luminosity model: this is the most complicated hypothesis to deal with. It assumes that the model accurately describes the phenomenon under analysis. However,

this assumption is highly unrealistic given the intricate nature of the phenomenon, which is influenced by numerous factors. Consequently, it is very difficult to develop a model that precisely captures the intricacies of this collision. This model is a simplification of reality, indeed does not incorporate some known aspects. However, it has the potential for extension and further development.

In this section, we will present some results based on the aforementioned hypotheses. Subsequently, an attempt was made to eliminate at least the first hypothesis in order to not only consider a more realistic scenario, but also to assess the sensitivity of the approach that will be presented when incorporating realistic random errors on the luminosity value.

The upcoming presentation will present the results in ascending order of complexity, first with the 1D (all equal emittances) scenario which is more straightforward to visualize, and culminating with the 4D (four distinct emittances) scenario which is inherently impossible to fully visualize due to its dimension.

The accelerator's initial configuration, specifically the starting values of its parameters, has been established based on the customary configuration system, as shown in the Table (5.1).

	parameters	value
$p_{cost}^{\rightarrow}$	f_{rev}	11 245 Hz
	n_b	2736
	N1,N2	1.4e11
	dispersion	0
	tot energy	6800 TeV
$p_{var}^{\rightarrow}$	$\Delta\mu_{0x}, \Delta\mu_{0y}$	0
	$\Delta\theta_{0x}$	0 μ rad
	$\Delta\theta_{0y}$	320 μ rad
	$\beta_{x1}, \beta_{x2}, \beta_{y1}, \beta_{y2}$	30 cm
	$\alpha_{x1}, \alpha_{x2}, \alpha_{y1}, \alpha_{y2}$	0 cm
	σ_{z1}, σ_{z2}	35 cm

Table 5.1: Table representing the different parameters to determine the initial configuration system.

where the notation Δp means $\Delta p = p^{B1} - p^{B2}$. For the purpose of simplicity, it has been assumed that there is no dispersion, leading to the conclusion that the transverse standard deviation of the beam profile is solely influenced by the betatronic component.

5.2.1 1D case : analytical solution

The present scenario is the only one that is well conditioned in advance. The objective is to determine a singular emittance ($\epsilon = \epsilon_{x1} = \epsilon_{x2} = \epsilon_{y1} = \epsilon_{y2}$) that is the same across all planes and beams, from a single equation. This equation correspond to the computation of luminosity within the initial configuration system. In this particular scenario, the computation of the model becomes more feasible due to the ability to calculate the integral presented in equation (4.10) along the longitudinal axis using analytical methods, resulting in a closed-form solution. Indeed, in the scenario where the configuration system is assumed to have a perfect head-on collision (with no crossing angle), that there is no significant dependence of the transverse axis on the longitudinal one, and it is possible also to assume that the relativistic beta value for both beams is identical and that the optics remain constant, $\beta(z) = \beta^*$. Under these conditions, the equation can be reformulated as follows:

$$L = \frac{N_1 N_2 n_b f_r}{\sqrt{2\pi^{\frac{3}{2}}} \sqrt{(\sigma_{x1}^2 + \sigma_{x2}^2)} \sqrt{(\sigma_{y1}^2 + \sigma_{y2}^2)}} \int_{-\infty}^{\infty} \frac{e^{\frac{-z^2 2\beta_r^2}{\beta_r^2 \sigma_{z1}^2 + \beta_r^2 \sigma_{z2}^2}}}{\sqrt{(\beta_r^2 \sigma_{z1}^2 + \beta_r^2 \sigma_{z2}^2)}} dz \quad (5.1)$$

Given that we are examining the scenario in which all the transverse emittances exhibit equal values, this can be reformulated as:

$$L = \frac{N_1 N_2 n_b f_r}{4\pi \beta_r \sigma^2} = \frac{N_1 N_2 n_b f_r}{4\pi \beta_r \beta^* \epsilon} \quad (5.2)$$

Hence, the emittance can be determined using analytical means by inverting the sole equation representing the luminosity value in the given reference scenario. Due to the well-conditioned nature of the problem, it is unnecessary to consider collisions from several perspectives, as the provided information is deemed adequate.

Furthermore, the underlying physics behind the structure can be deemed even more intricate. Specifically, if we imagine that the crossing angle deviates from zero in one of the two transverse axes (in the y-axis), it becomes feasible to employ the previously mentioned formula presented in the previous chapter (4.11) within equation (4.10). From this new luminosity formula, by maintaining the assumption that the relativistic beta is equivalent for both beams and that the optics remain constant, it remains possible to calculate a closed form solution, although it is more complex:

$$\begin{aligned} L &= N_1 N_2 n_b f_r \frac{e^{\frac{-\frac{\mu_{x1}^2}{2} + \mu_{x1} \mu_{x2} - \frac{\mu_{x2}^2}{2}}{2\sigma^2}}}{2\sqrt{2}\pi^{\frac{3}{2}}\sigma} \int_{-\infty}^{\infty} \frac{e^{\frac{-\frac{\mu_{y1}^2}{2} + \mu_{y1} \mu_{y2} - \frac{\mu_{y2}^2}{2}}{2\sigma^2}} e^{\frac{-z^2 2\beta_r^2}{\beta_r^2 \sigma_{z1}^2 + \beta_r^2 \sigma_{z2}^2}}}{\sigma \sqrt{(\beta_r^2 \sigma_{z1}^2 + \beta_r^2 \sigma_{z2}^2)}} dz \\ &= N_1 N_2 n_b f_r \frac{e^{\frac{-\frac{\mu_{x1}^2}{2} + \mu_{x1} \mu_{x2} - \frac{\mu_{x2}^2}{2}}{2\sigma^2}}}{2\sqrt{2}\pi^2\sigma} \frac{e^{\frac{-\frac{2(\mu_{y1} - \mu_{y2})^2}{4\sigma + (p_{y1} - p_{y2})^2(\sigma_{z1}^2 + \sigma_{z2}^2)}}}{\beta_r \sqrt{(\sigma_{z1}^2 + \sigma_{z2}^2)((p_{y1} - p_{y2})^2 + \frac{4\sigma^2}{\sigma_{z1}^2 + \sigma_{z2}^2})}} \\ &= N_1 N_2 n_b f_r \frac{e^{\frac{-\frac{\mu_{x1}^2}{2} + \mu_{x1} \mu_{x2} - \frac{\mu_{x2}^2}{2}}{2\beta\epsilon}}}{2\sqrt{2}\pi^2\sqrt{2}\beta\epsilon} \frac{e^{\frac{-\frac{2(\mu_{y1} - \mu_{y2})^2}{4\sqrt{\beta\epsilon} + (p_{y1} - p_{y2})^2(\sigma_{z1}^2 + \sigma_{z2}^2)}}}{\beta_r \sqrt{(\sigma_{z1}^2 + \sigma_{z2}^2)((p_{y1} - p_{y2})^2 + \frac{4\beta\epsilon}{\sigma_{z1}^2 + \sigma_{z2}^2})}} \end{aligned} \quad (5.3)$$

Finally, it is worth noting that a closed form for luminosity calculation exists. However, determining the emittance from this closed form is currently not a straightforward task.

5.2.2 2D cases

In this scenario, the objective is to increase the degrees of freedom in order to obtain two distinct emittances within the beams or planes. In order to condition the problem, it is necessary to introduce an additional equation.

There is a preference in choosing which parameter to shift in order to obtain a different perspective of the collision. Indeed, as already indicated, certain parameters, especially β^* and α , pose challenges in terms of manipulation. Shifting these parameters is a non-trivial task, requiring a complex technique, and accurately measuring their values is not straightforward.

Among the various parameters that can be adjusted, there exist certain ones that are relatively more straightforward to manipulate. Specifically, these parameters pertain to the orbit of the beams, and they encompass the offset and the crossing angle of the two beams.

Another easily modifiable parameter that can be accurately measured is the standard deviation of the longitudinal beam profile. This choice raises the question on whether altering the longitudinal beam profile to manipulate luminosity values is truly indicative of the system and can provide insights into the transverse beam profile. The answer is affirmative, as the problem exhibits a strong coupling, so adjusting the longitudinal beam profile is informative in order to effectively invert the problem.

In theory, it is feasible to use all of these parameters inside our configuration system to achieve a problem inversion. However, the parameters most commonly employed are those associated with the orbit of the beam.

Given that the problem is two-dimensional, it is sufficient to have only one additional equation in addition to the one defined by the configuration system. Therefore, only one parameter needs to be shifted. Specifically, the following results will be presented by shifting only the offset on the x-axis (i.e., $\mu_{x1} = \mu_{x1,0} + c/2$, $\mu_{x2} = \mu_{x2,0} - c/2$). The selection of this parameter is made arbitrarily, as the outcomes remain consistent regardless of the specific parameter picked among those associated with the orbit. To demonstrate this, the final result will be shown for all the orbit parameters.

In all the various numerical analysis, unless explicitly stated otherwise, we will use the beta function in its whole complexity as expressed by equation (4.6).

Different emittances in the two transverse axis

In this case, let us consider the hypothetical scenario where the emittances between the two beams along the same axis are identical ($\epsilon_x = \epsilon_{x1} = \epsilon_{x2}$, $\epsilon_y = \epsilon_{y1} = \epsilon_{y2}$). We will use the initial conventional configuration system with a non-zero plane crossing angle on the y-axis, and afterwards with an offset shift on the x-axis, as previously depicted. By employing the LS, it becomes feasible to attain the desired emittances. Figure (5.1) depicts the iteration stack of the LS in order to attain the desired emittance on the system's penalty function.

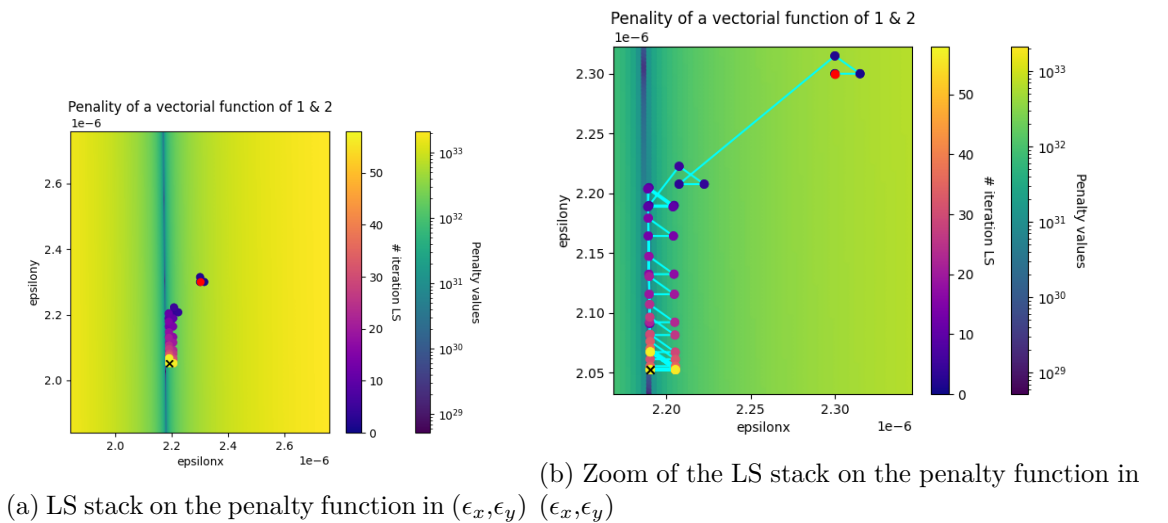


Figure 5.1: At left, the representation of the stack in the square outlining the 20% of the first estimation, at right, the zoom of the same plot on left.

The initial observation pertains to the distinctive structure of the penalty function, which

exhibits a rift-like pattern similar to a line where all values converge towards zero. However, this perception is misleading and arises from the differences in magnitude, as there exists a substantial discontinuity within a limited range. This is evident in Figure (5.2), which represents the penalty function values within this rift, clearly indicating the presence of a global minimum along that. The observed rift can be explained by the non-zero crossing angle of the y-axis. The penalty function used in this analysis is based on the least squares method, where the two equations are obtained by a subtraction of a function and a scalar value. Therefore, the dependence illustrated in Figure (5.2) is solely determined by the luminosity function. The observed dependence is not unexpected, as it indicates that the luminosity model exhibits greater sensitivity to variations in emittance within the plane where there is no crossing angle; so where the two beams are perfectly aligned, and it is logical that the collision in that axis would be more sensitive.

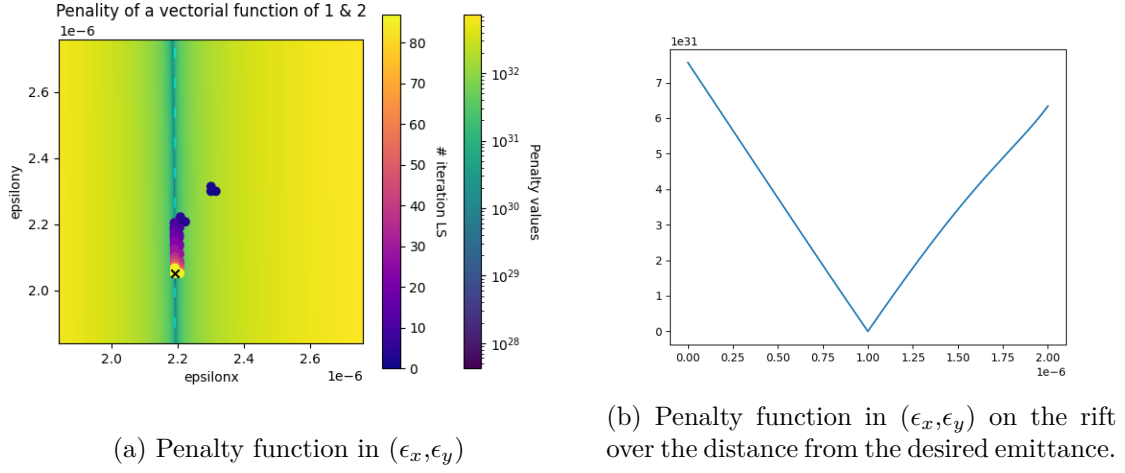


Figure 5.2: At left, the representation of the penalty function in the 2D plane with the representation of the different LS iterations, at right, the plot of the penalty function in the dashed cyan line of the plot on the left.

Another notable aspect to observe is the configuration of the stack, which clearly demonstrates the presence of the intrinsic procedure and the local approximation technique employed to attain the global minimum, as described in (5.1.2). In the last iterations, as the rift has been achieved, the penalty value of the function exhibits a notably tiny magnitude. To emphasise this characteristic, the relative errors of the luminosity estimation in the last 20 iterations of the LS have been graphically represented in Figure (5.3). The plot clearly illustrates the small order of magnitude of the relative error in the last iterations, despite their imprecise estimation of epsilon.

To ensure the validity of the approach, the procedure was replicated for a total of two hundred distinct selections of (ϵ_x, ϵ_y) . Each component was chosen randomly within a range of 20% relative to the initial estimate of $2.3e - 6$. Figure (5.4) depicts the precision of this approach, considering shifting approach not only in μ_x , but also in μ_y , θ_x , and θ_y .

The Table (5.2) displays the average relative error and average calculation time for the least squares method across all 200 emittances, for all potential orbit shifts.

For the purpose of clarity, the subsequent sections will solely present the final figure, unless there are unusual cases that require otherwise, representing the accuracy of the techniques across a total of two hundred cases.

Relative error on Luminosity estimation over the last 20 iterations

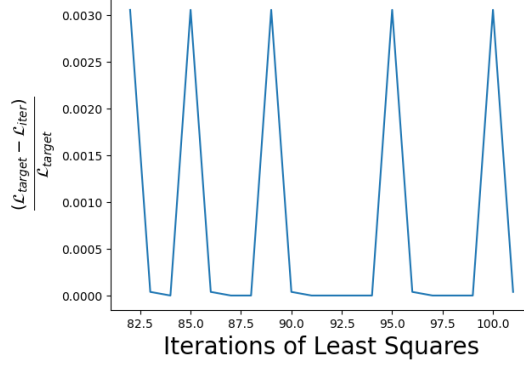
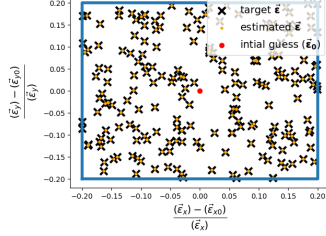
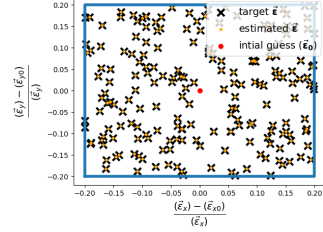


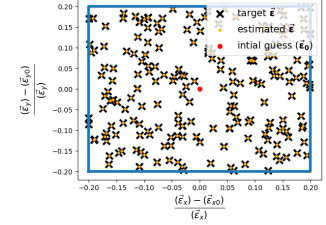
Figure 5.3: Relative errors on the luminosity value of the last 20 iterations of the LS, in order to reach the best solution

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the offset on y (μ_x)


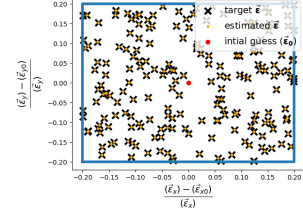
(a) Shift for the offset in the x-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the offset on y (μ_y)


(b) Shift for the offset in the y-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying crossing angle on x (θ_x)


(c) Shift for the crossing angle in the x-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the crossing angle on y (θ_y)


(d) Shift for the crossing angle in the y-axis

 Figure 5.4: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_x, \epsilon_y)$ chosen randomly in the cyan square of 20% with respect to the initial guess, with different shifting parameter

Different emittances for the two beams

In this scenario, let us consider the equal emittances between the two axes within the same beam ($\epsilon_1 = \epsilon_{x1} = \epsilon_{y1}, \epsilon_2 = \epsilon_{x2} = \epsilon_{y2}$). We will employ the conventional configuration system with a non-zero plane crossing angle on the y-axis, and afterwards with an offset shift on the x-axis, as previously depicted. It is possible to see that, in comparison to the previous case, it is not feasible to attain the desired emittance using the LS. Indeed, the iteration stack of the LS is represented on the system's penalty function, in the Figure (5.5), and it is evident that it does not achieve the desired solution.

The analysis of the image reveals two distinct observations. Firstly, the LS successfully reaches the rift, but experiences troubles in finding the desired emittance. Additionally, it is

	(ϵ_x, ϵ_y) without err	LS time
μ_x	1.1658782597902616e-12	0.11228418350219727
μ_y	9.164000124614319e-14	0.04223036766052246
θ_x	9.041838560737559e-15	0.0397799015045166
θ_y	3.892071327841851e-13	0.06345224380493164

Table 5.2: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, for all the possible shifts of the orbit.

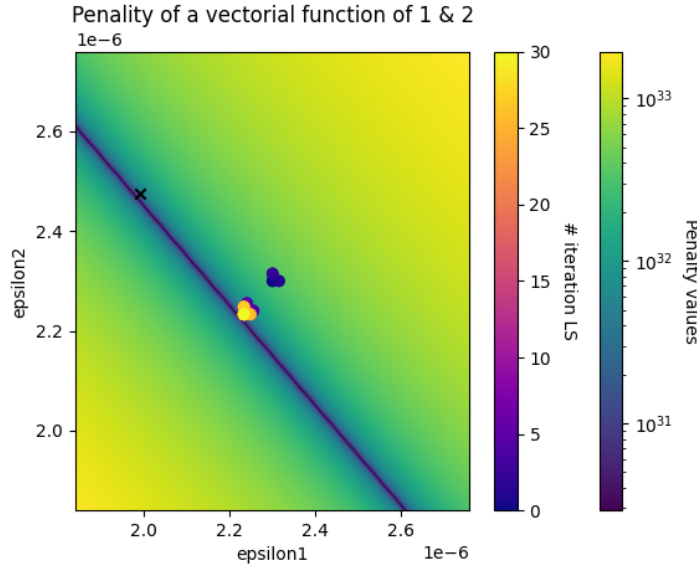
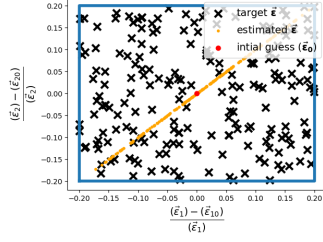


Figure 5.5: Relative errors on the luminosity value of the last 20 iterations of the LS, in order to reach the best solution

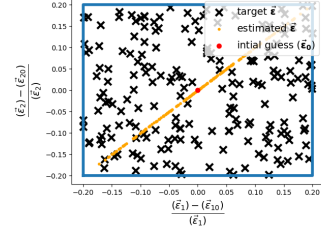
observed that the rift occupies an anti-diagonal position within the two-dimensional plane. Consequently, the sum of the LS result's components results is identical to the sum of the desired emittances, indicating a potential symmetry issue. This deduction aligns with the visual representation provided in Figure (5.6). From the computation of the penalty function on the rift is illustrated in Figure (5.7), it can be observed that the values on the rift are mainly null. Consequently, despite the addition of another equation, the problem remains ill conditioned. This result, from the analysis of the luminosity model, isn't so surprising.

From a technical point of view, the primary objective is to achieve distinct emittances for each beam. However, the selected initial configuration system for the model is entirely symmetrical, respecting this requirement across all parameters. Hence, the objective is to discern a distinction between the two beams in the unknown variables, while assuming all other parameters to be perfectly symmetrical. By symmetrically manipulating a parameter that is inherently symmetrical, it was predictable that this action would not provide any meaningful outcome. It is necessary to disrupt the symmetry inside the configuration system. However, a pertinent question arises whether employing a single random parameter is adequate for achieving this objective.

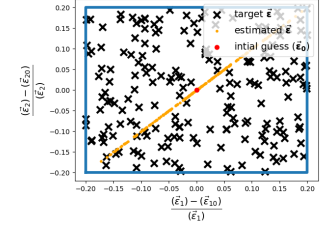
The model reveals a visible relation with certain parameters (β, ϵ) , given that the standard deviation of off-momentum is zero

Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_x)


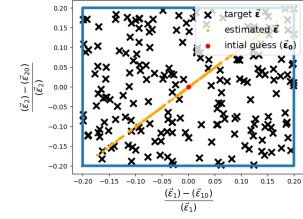
(a) Shift for the offset in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_y)


(b) Shift for the offset in the y-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying crossing angle on x (θ_x)


(c) Shift for the crossing angle in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the crossing angle on y (θ_y)


(d) Shift for the crossing angle in the y-axis

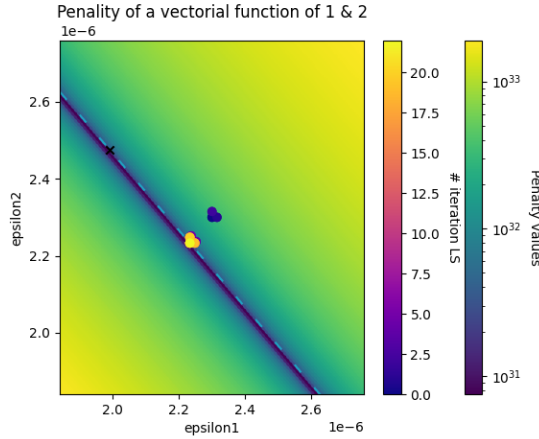
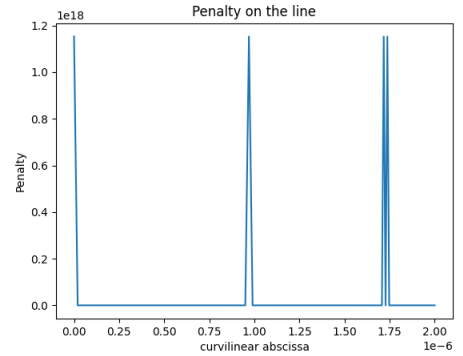
 Figure 5.6: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_1, \epsilon_2)$ chosen randomly in the cyan square of 20% with respect to the initial guess, with different shifting parameter

 (a) Penalty function in (ϵ_1, ϵ_2)

 (b) Penalty function in (ϵ_1, ϵ_2) on the rift over the distance from the desired emittance.

Figure 5.7: At left, the representation of the penalty function in the 2D plane with the representation of the different LS iteration, at right, the plot of the penalty function in the dashed cyan line of the plot on the left.

$$\mathcal{L}(\vec{\epsilon}, \vec{\beta} | \dots) = N_1 N_2 n_b f_r \int_{-\infty}^{\infty} \frac{e^{-\frac{\mu_x^2}{2} + \mu_{x1} \mu_{x2} - \frac{\mu_{x2}^2}{2}}}{\beta(\epsilon_1 + \epsilon_2)} \frac{e^{-\frac{\mu_y^2}{2} + \mu_{y1} \mu_{y2} - \frac{\mu_{y2}^2}{2}}}{\beta(\epsilon_1 + \epsilon_2)} e^{\frac{z^2 \left(-\frac{\beta_{r1}^2}{2} + \beta_{r1} \beta_{r2} - \frac{\beta_{r2}^2}{2} \right)}{\beta_{r1}^2 \sigma_{z1}^2 + \beta_{r2}^2 \sigma_{z2}^2}} \frac{1}{\sqrt{2\pi}^{\frac{3}{2}} \sqrt{\beta(\epsilon_1 + \epsilon_2)} \sqrt{\beta(\epsilon_1 + \epsilon_2)} \sqrt{(\beta_{r1}^2 \sigma_{z1}^2 + \beta_{r2}^2 \sigma_{z2}^2)}} dz = \mathcal{L}(\beta(\epsilon_1 + \epsilon_2) | \dots)$$

This observation highlights the fact that regardless of the parameter selected to introduce

asymmetry in the configuration system until the beta function is equal for both beams, the least squares method always results in identical outcomes. This is due to the LS method's inability to discern any distinction in that particular dependence, resulting in estimating the sum and producing the same value for both components. Therefore, it is necessary to disrupt the symmetry into the configuration system by altering either the β^* or the α^* .

From a practical point of view, this alteration presents a considerable level of complexity due to the uncommon requirement of modifying the optical properties of the machine in a distinct manner for B1 as compared to B2. This is primarily due to the fact that the optical characteristics are inherent to the machine itself, rather than being specific to the beam. Consequently, this request appears almost impractical from a technical standpoint. However, in order to facilitate the LS in achieving the desired emittance, it is not necessary to make significant modifications of the Twiss parameters. Specifically, a change on the order of $1e-4$ in the β^* is required, which is smaller than $1e-3$ of the original value chosen in the Table (5.1). Hence, it is necessary to establish that $\beta_1^* = \beta^*$ and $\beta_2^* = \beta^* + 1e-4$. The justification for this modification in the value of β^* can be readily explained by noting that the relative error associated to the measurement of this parameter exceeds $1e-3$. Consequently, this discrepancy in the two beta functions is already a component of the problem.

By implementing this slight shift in order to break the symmetry between the two beams, the outcomes face a full transformation, resulting in the successful attainment of the desired emittance, as illustrated in Figure (5.8) and Figure (5.9)

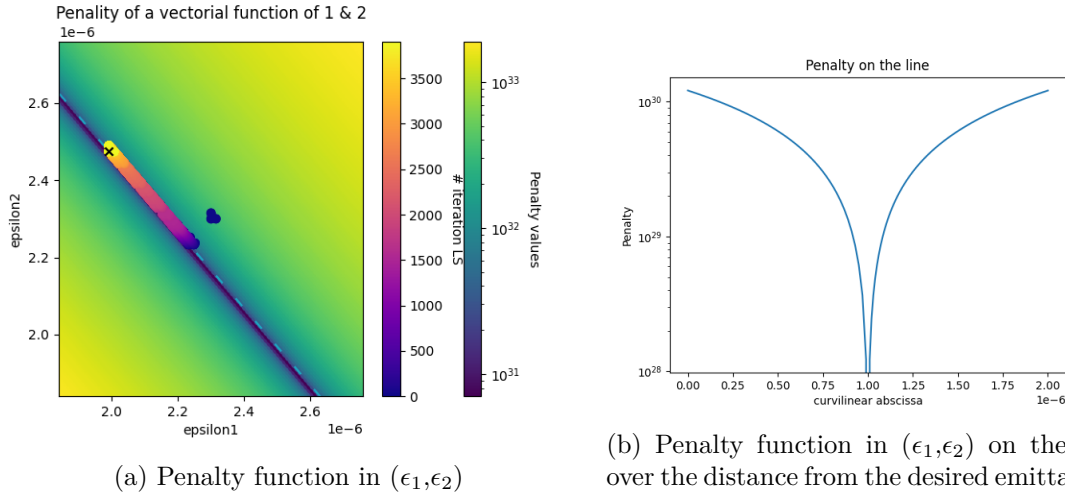
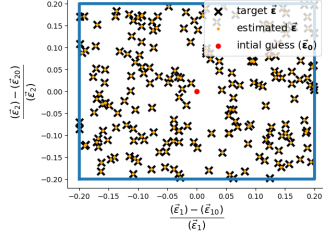


Figure 5.8: At left, the representation of the penalty function in the 2D plane with the representation of the different LS iteration, at right, the plot of the penalty function in the dashed cyan line of the plot on the left. Both cases when the configuration system in which the beta functions between the two beams are different ($\beta_1^* = \beta^*$ and $\beta_2^* = \beta^* + 1e-4$).

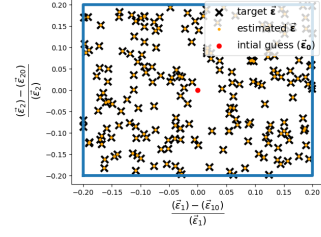
The Table (5.3) displays the average relative error and average calculation time for the least squares method across all 200 emittances, for all potential orbit shifts.

5.2.3 4D case: realistic case

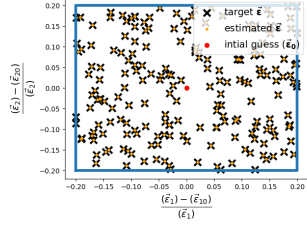
In this particular scenario, all the emittances have distinct values, this renders the observation of the penalty function's behaviour across the unknown variables more complex. This is due to the fact that the function would be plotted within a five-dimensional space. Therefore, when necessary, only a partial representation of this function will be depicted, focusing

Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_x)


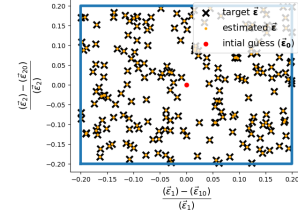
(a) Shift for the offset in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_y)


(b) Shift for the offset in the y-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying crossing angle on x (θ_x)


(c) Shift for the crossing angle in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the crossing angle on y (θ_y)


(d) Shift for the crossing angle in the y-axis

Figure 5.9: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_1, \epsilon_2)$ chosen randomly in the cyan square of 20% with respect to the initial guess, with a configuration system in which the beta functions between the two beams are different ($\beta_1^* = \beta^*$ and $\beta_2^* = \beta^* + 1e - 4$), with different shifting parameter

	(ϵ_1, ϵ_2) without err	LS time
μ_x	4.3193644987350564e-09	3.666034460067749
μ_y	1.313205239961527e-10	0.6698188781738281
θ_x	1.8395420483684256e-07	5.77882194519043
θ_y	1.6103278738191488e-07	5.784540414810181

Table 5.3: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, for all the possible shifts of the orbit.

on the projection around our initial guess of the least squares.

The objective in this case is to apply the knowledge gained from prior cases to achieve a favourable outcome in this most general scenario. Therefore, based on the previous case analysis, the initial configuration system employed is the one where the β^* functions vary between the beams. Given the presence of four unknown variables in this particular scenario, it is insufficient to shift a single parameter. Indeed, it can be easily demonstrated that only adjusting the offset along the x-axis will not yield the intended outcome. From a mathematical perspective, it is particularly predictable due to our objective of determining four distinct unknowns from a system consisting of only two equations leading to an under-determined system. Hence, in order to achieve the desired solution, it is necessary to have a system consisting of a minimum of four well-conditioned equations. In order to enhance computational efficiency and facilitate notation, we will use a system of five equations. This system will include shifting for both the crossing angle and the offset for both axes. However, empirical testing has demonstrated that identical outcomes may be attained by adjusting only three out of the four aforementioned parameters. In this particular case, it is intriguing to examine two distinct configurations of the machine.

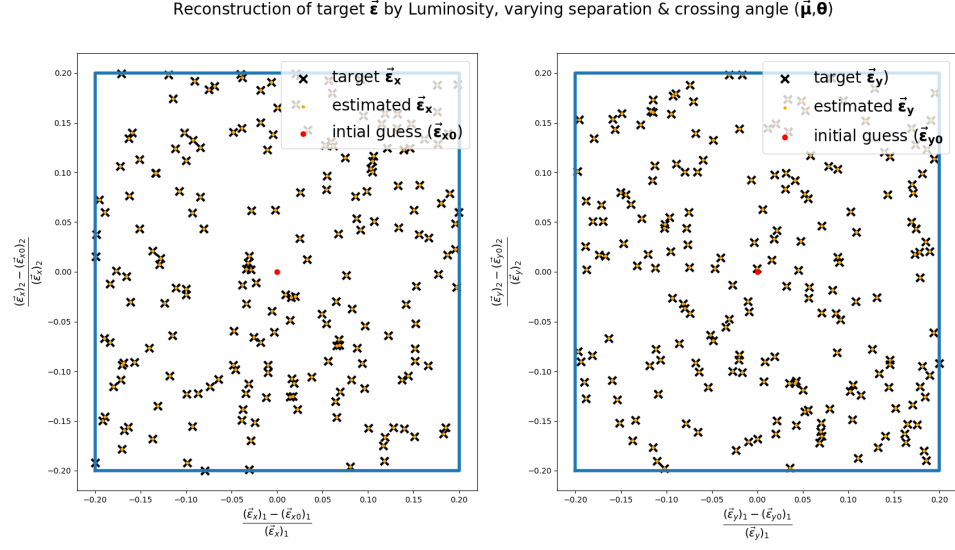


Figure 5.11: Relative errors on the luminosity value of the last 20 iterations of the LS, in order to reach the best solution

	$(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ without err	LS time
orbit	1.2571219944678352e-09	5.893694162368774

Table 5.4: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non linear LS estimation, obtained shifting the different parameters of the orbit

were solely achieved by shifting the orbit (crossing angle and offset). However, it is worth noting that modifying the optics (β^* and α^*), would yield similar outcomes. The only two exceptions to this are: the 2D scenario where the two beams possess different emittances with the initial configuration system with identical optics and this particular scenario. In the latter case, when the orbit is shifted, the LS is able to get the desired emittances, as demonstrated in Figure (5.12). However, when the optics are shifted, as shown in Figure (5.13), the LS may encounter certain difficulties. This result, considering the physics behind the collision, is not so surprising.

In particular, the Twiss parameters have been formally defined as optical functions that are associated with the magnets along the accelerator ring. Given that one can be derived from the other two, our analysis will solely concentrate on the parameters β^* and α^* . The first one pertains to the focusing or defocusing capability of the quadrupole magnets. Indeed, as depicted in Figure (4.6), decreasing the value of β^* increases the focusing ability of the quadrupole magnets. On the other hand, the second parameter is primarily associated with the position of the minimum, illustrated in Figure (4.6). When it deviates from zero, the origin of the parabola shifts in relation to the interaction point. Based on the concise explanation provided regarding the parameters, it is not surprising that the desired emittance cannot be achieved through the LS shifting of optics in this particular scenario. This is primarily due to the absence of a crossing angle, wherein the bunches are already optimally aligned. Consequently, introducing a new configuration and equation shifting the optics function would not yield additional insights into the overall system, resulting in an under-conditioned problem.

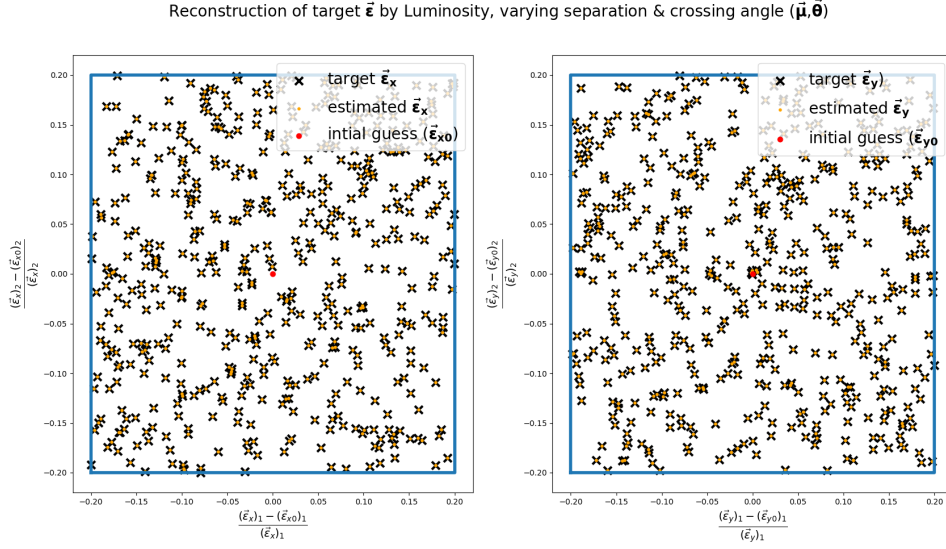


Figure 5.12: Relative errors on the luminosity value of the last 20 iterations of the LS, in order to reach the best solution

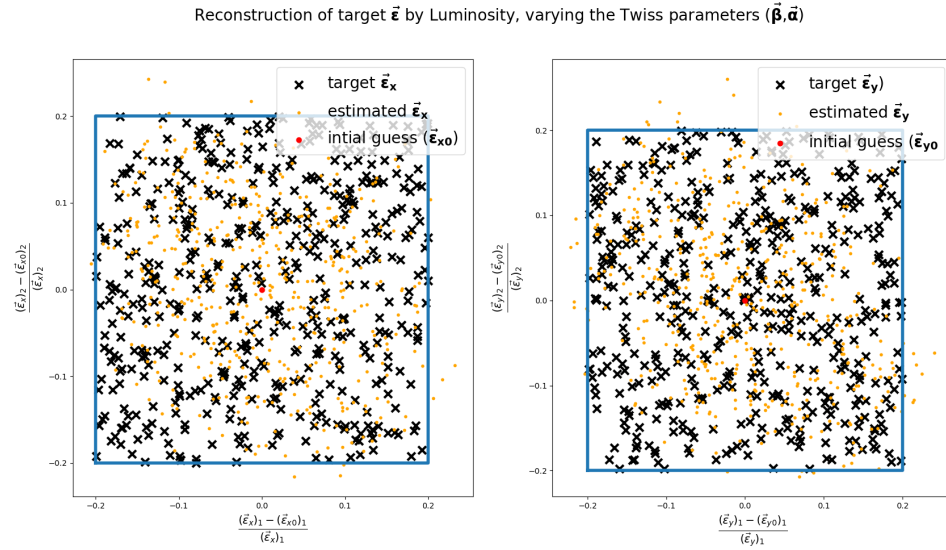


Figure 5.13: Relative errors on the luminosity value of the last 20 iterations of the LS, in order to reach the best solution

5.3 2nd inversion problem, with random errors in the Luminosity output

The reason underlying this approach aligns with the introduction of Section (5.2). The ideal setup of the luminosity formula is as follows:

$$\mathcal{L} = \mathcal{L}(\vec{\epsilon} | p_{var}, p_{const})$$

Shifting the variable p_{var} , we aim to derive a set of equations that can be minimized using the Non-linear Least Square method. In the previous section, the comparison between

this simulation and a practical approach on the accelerator has revealed the presence of some assumptions made by the simulation that may not be immediately apparent. The purpose of this section is to eliminate one of these hypotheses, specifically the assumption that the value of luminosity is accurately measured. Instead, it will be assumed that there is a random error associated with its measurement. This generalization slightly modifies the underlying structure of the simulation:

- From the selected initial configuration system for a detector, a decision has been made on which parameter to modify, resulting in a new value for that parameter.
- Selecting the four emittances randomly, which are assumed to be unknown. These emittances are then provided as input to the model with the objective of calculating the luminosity ($\mathcal{L}$).
- Obtained the luminosity $\mathcal{L}$, **it is combined with a random error derived from a random sample following a normal distribution with mean $\mu = 0$ and standard deviation $\sigma = \mathcal{L} * 1e - 3$.**
- **The value of the luminosity with random error, $\mathcal{L} + \Delta\mathcal{L}$ it is used in the equation, by subtracting it from the model with the chosen configuration (the one that when subjected to randomly chosen emittances returns $\mathcal{L}$).**
- The equation is incorporated into the system, subsequently it has been decided another parameter for shifting, and the process restarts from the original stage.

The simulation has resulted in the development of a system wherein each equation exhibits a diverse sample error in luminosity value compared to the subtracted model.

The idea of this section is to test the proposed approach, as the previous section revealed fascinating outcomes with a nearly perfect estimation of the desired emittances. The objective of this section is to evaluate the sensitivity of this approach to random errors in luminosity, with the aim of determining the impact on the estimation of emittances. This error pertaining to luminosity, given the complexity of its measurement, is notably minimal. This is the primary reason why we have chosen to use luminosity as the central element of this inversion.

Given the recent inclusion of a random error in the luminosity model, the task of using the luminosity formula and the underlying physics to determine the cause for the different outcomes in the LS has become more challenging. Due to these reasons, this particular phase of the project will prioritise a “numerical” approach rather than exploring into the physics behind. Indeed, it seeks to use the penalty function’s representation in many domains to address and overcome various challenges.

To evaluate the impact of random errors on luminosity in the simulation, the first two examples will be examined without modifying the approach, in order to observe how the estimation is affected.

5.3.1 Adding the errors in the approach used before: 1D case

The initial situation under examination pertains to the well-defined scenario of a single equation for a solitary unknown. It is assumed that the emittances, both across planes and beams, are identical ($\epsilon = \epsilon_{x1} = \epsilon_{x2} = \epsilon_{y1} = \epsilon_{y2}$). Based on the preceding section, it has been established that the integral inside the luminosity formula can be evaluated in a closed form. Therefore, it was unnecessary to employ the least squares method to address the problem in question, as it may be resolved using analytical means. From this well-definition of the

problem, it has been expected that this particular scenario will exhibit minimal sensitivity to random errors in luminosity. Indeed, the sensitivity is expected to be exclusively determined by the dependence on emittance in the luminosity closed form. In the present context, it is possible to resolve the equation with analytical techniques. Nevertheless, in order to maintain consistency with the other cases, also this has been analysed by employing the least squares method. This method enables a computational analysis of the influence on the penalty function. The penalty function for different values of ϵ is depicted in Figure (5.14), displaying both the presence and absence of error in the same scenario. The use of this graphic depiction facilitates a comprehensive comparison of the two cases.

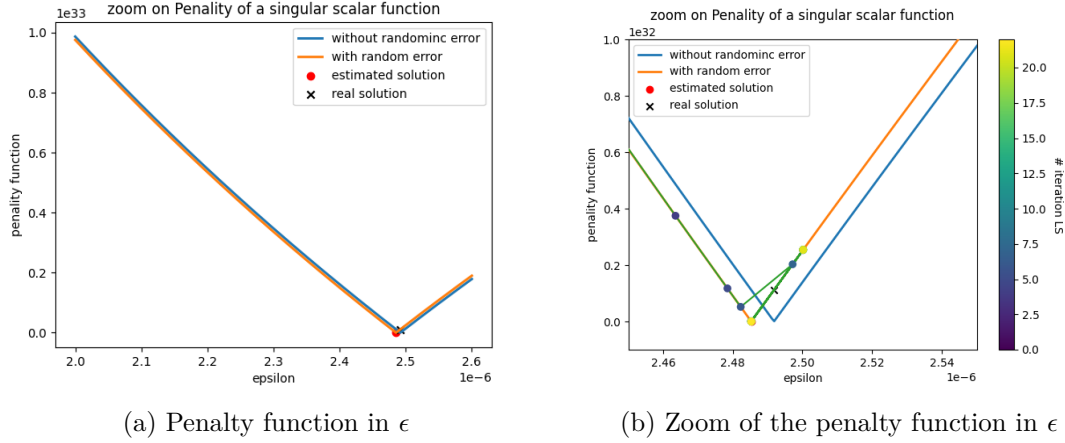


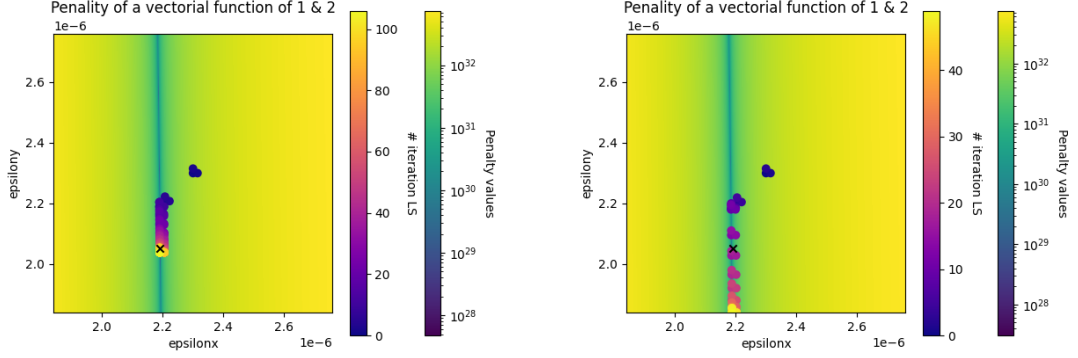
Figure 5.14: At left, the representation of the penalty functions with and without error, with underlined the respective minimum, at right the zoom of the functions on left near the minimum with the representation of the different iteration of the LS on the penalty function with error.

The Figure (5.14b) illustrates that the penalty function with error in the 1D case displays a noticeable displacement from the ideal function. This displacement has the potential to result in an inaccurate assessment of the ideal ϵ . The estimation exhibits a relative inaccuracy of approximately 2%, that is particularly promising given that the current instruments employed for emittance estimation in the LHC possess a 10% margin of relative error.

5.3.2 2D cases

The second case presents a more intricate challenge, as it is no longer straightforward to assert that the problem is well-defined. In this scenario, we are assuming that the emittances are the same for the beams but differ on the planes, ($\epsilon_x = \epsilon_{x1} = \epsilon_{x2}$ and $\epsilon_y = \epsilon_{y1} = \epsilon_{y2}$). In this particular scenario, the integral inside the luminosity formula hasn't a closed form, hence rendering it impossible to analytically compute the inversion of the luminosity. Based on the complexity of this situation, it is expected that the relative inaccuracy of the emittance estimation is at least equal to, if not bigger than, that of the prior case. In each of the several figures presented within this subsection, unless otherwise specified, the desired outcome will be attained by shifting the offset on the x-axis for both beams, as illustrated in the Section (5.2.2). The Figure (5.15) always illustrates the same approach with two different random error on the luminosity, properly to visualize the sensitivity of the approach defined.

The image clearly demonstrates that when the order of the error increases towards our desired value, the LS method is notably misleading in estimating the desired emittance. The least squares method is able to accurately identify the rift, approaching the desired value



(a) Penalty function in (ϵ_x, ϵ_y) , with random error sample from generated by a normal distribution with $\sigma = \mathcal{L} * 1e - 6$
 (b) Penalty function in (ϵ_x, ϵ_y) , with random error sample from generated by a normal distribution with $\sigma = \mathcal{L} * 1e - 3$

Figure 5.15: At left, the representation of the penalty functions with error of the order $1e-6$, with underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value, at right the representation of the penalty functions with error of the order $1e-3$, with underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value. Both the approach has been done in a configuration system with plane crossing angle on the y-axis.

ϵ_x with an error of 2%. However, when it comes to the deviation along the rift, the LS method fails to achieve the desired value ϵ_y and reaches the boundary of the LS, leading to a total relative error, for the entire vector, of 20%. The present error is similar to the one of the employed devices within the accelerator. However, our primary objective is to enhance the accuracy of the emittance estimation.

Based on this concise study, it can be asserted that even in the two-dimensional scenario with error, it is possible to observe the inherent pathological structure of the penalty function. The aforementioned misleading estimation is not directly relevant to the problem at hand, as it remains well-conditioned. However, the approach is significantly affected by a strong sensitivity to random errors in luminosity, which in turn results in an undesired unique global minimum of the penalty function.

5.3.3 Approaches to overpass this sensitivity

The objective of this section is to use the penalty function's structure to develop an approach for effectively mitigating this sensitivity.

One potential strategy to consider is to systematically increase the quantity of samples within the selected configuration. When selecting a configuration system for a specific detector, it is important to note that in the case of a circular collider, it is feasible to achieve consistent measurements of luminosity through numerous repetitions, each with a different random error. This is primarily due to the high velocity of the beam, which approaches the speed of light. This would result, for the law of large numbers, in a reduction of the relative error associated with emittance estimation. This is primarily due to the mitigation of the effect of the random errors in luminosity on the overall penalty function. Unfortunately, this approach is not feasible due to our objective of achieving the appropriate emittance values in the shortest possible time. Waiting for extended periods in a particular configuration to acquire diverse data is a luxury that is beyond our capability. The optimal strategy

should establish a connection between the objectives of acquiring additional measurements in a single configuration system and of accelerating the transition between the different configuration systems.

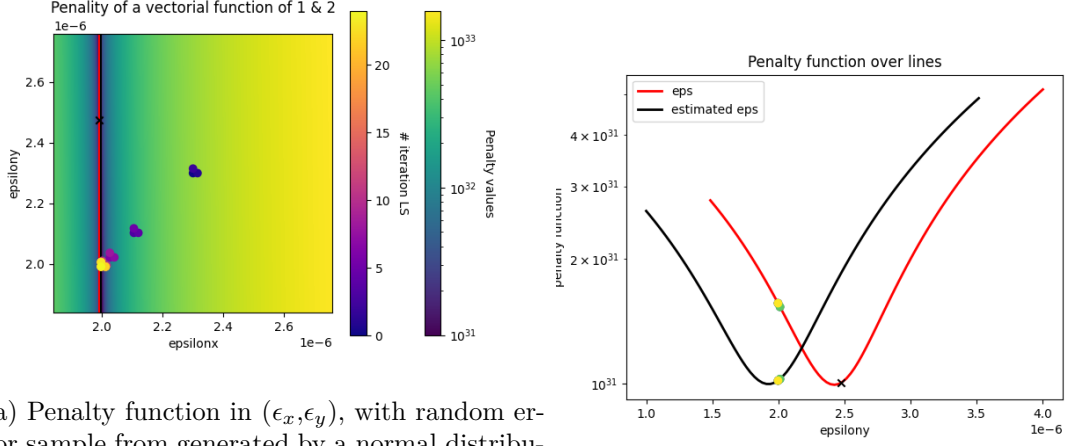
The attainment of this ideal strategy may appear unachievable due to the simultaneous pursuit of two conflicting objectives, but this argument is not entirely accurate. The transition between two configuration systems is achieved through a constant shift in a detector that observes bunch collisions occurring on the order of nanoseconds. The detector, prior to reaching the desired configuration system, observes collisions between bunches in numerous potential configurations. Unfortunately, a full review of thousands of configurations, before arriving to the desired one, is unfeasible due to the prevalence of rumours and errors below a certain threshold. Consequently, the effort would yield no meaningful results. Given the presence of a random error on the luminosity with a maximum magnitude of $1e - 3$, and taking into account that the desired shifts in the parameter would result in an impact on the luminosity value of 1%, the objective is to select just 10 configurations from the numerous available options. Each of these configurations would induce a luminosity alteration of 1%. By adopting this technique, it becomes feasible to increase the quantity of luminosity measurement, so moving closer to the objective of mitigating the impact of random errors on the final penalty function, without affecting the transition between configurations system. Unfortunately, as will be elucidated afterwards, this particular strategy, without further techniques, is insufficient to achieve the desired emittance.

Different emittances in the two transverse axis

In this particular case, we are examining the identical scenario as previously discussed ($\epsilon_x = \epsilon_{x1} = \epsilon_{x2}$, $\epsilon_y = \epsilon_{y1} = \epsilon_{y2}$), with a plane crossing angle on the y-axis. Furthermore, instead of solely focusing on the shifted configuration system and the initial one, we are incorporating other nine configurations that can be precisely measured within the transition, based on the insights gained from the prior analysis. This slightly modifies the underlying structure of the simulation:

- From the **selected configuration system** for a detector, a decision has been made on which parameter **to change by 1%the luminosityin relation to the configuration system**, resulting in a new value for that parameter.
- Selecting the four emittances randomly, which are assumed to be unknown. These emittances are then provided as input to the model with the objective of calculating the luminosity ($\mathcal{L}$).
- Obtained the luminosity $\mathcal{L}$, it is combined with a random error derived from a random sample following a normal distribution with mean $\mu = 0$ and standard deviation $\sigma = \mathcal{L} * 1e - 3$.
- The value of the luminosity with random error, $\mathcal{L} + \Delta\mathcal{L}$ it is used in the equation, by subtracting it from the model with the chosen configuration (the one that when subjected to randomly chosen emittances returns $\mathcal{L}$).
- This equation is added in the system, afterwards **using the same parameter with the same value, and then resetting the process**.

Unfortunately, using this approach is insufficient, as the system fails to entirely eliminate the luminosity error over the penalty function. Consequently, the LS method does not yield the desired emittance, as depicted in Figure (5.16).



(a) Penalty function in (ϵ_x, ϵ_y) , with random error sample from generated by a normal distribution with $\sigma = \mathcal{L}^* 1e - 3$, with two vertical lines that has the abscissa of the estimated emittance and of the true emittance (b) Penalty function in ϵ_y , at ϵ_x fixed the minimum found with the LS is smaller than the desired emittance

Figure 5.16: At left, the representation of the penalty functions with error of the order $1e-3$, with underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value, and with two vertical lines that has the abscissa of the estimated emittance and of the true emittance. At right the representation of the penalty functions over the ϵ_y of the two lines of the plot on left, with underlined the ideal value of emittance and the different iteration, on both the functions, of the LS in order to reach the goal value. Also, if complicated to notice, the black function is smaller than the red one in order of $1e29$.

In addition to the aforementioned undesirable outcome, it is interesting to observe another significant aspect. The difference between the penalty function's structure in (5.16) and the one depicted in Figure (5.15) is apparent. Specifically, the rift in Figure (5.16a) is noticeably wider compared to the preceding figure, and it contains values that are not so small. The observed distinction is not solely attributable to the random selection of errors or goal emittances, but rather is mostly influenced by the chosen approach. Indeed, remarking the structure of the LS, as shown in the subsection (5.1.2), the overall penalty function is obtained by squaring the sum of different 11 functions that share the same structure with the rift. However, considering that each of the 11 equations in the system contains a distinct random error, this rift slightly shifts its position along the x-axis, but particularly affecting the location of the minimum along the y-axis. The reason for the increased thickness of the final penalty function's rift, as well as the larger values it contains, can be attributed to this factor. Additionally, due to the same underlying factor, the penalty function exhibited by the two lines within the rift, as depicted in Figure (5.16b), demonstrates a smoother behaviour compared to previously illustrated functions.

Hence, the proposed approach is insufficient in achieving the target emittance mostly due to the intrinsic structure of the penalty function, which exhibits a rift along the y-axis. Given the reasoning behind this rift exposed in (5.2.2) and the underlying concept of the LS method in the SciPy library (5.1.2), it may be interesting to explore the possibility of mitigating the aforementioned rift phenomenon by using the data from two detectors simultaneously, rather than relying solely on the information from a single detector. Specifically, within the Large Hadron Collider, the two detectors that generate the highest luminosity are ATLAS and CMS. These detectors experience the same collisions between bunches of particles, with the main difference being that the planes in which the crossing angle is non-zero are precisely opposite to each other. Changing the plane for crossing angle results in a different penalty

function, as depicted in Figure (5.17a). It is worth noting that while the specific value of this angle in each detector may vary, the orientation of the planes remains opposite, with one corresponded along the x-axis and the other along the y-axis.

Based on the assumption that the emittance of the beam remains constant during its traversal of the pipe, it is feasible to employ two distinct initial configuration systems. One system involves a non-zero crossing angle in the y-axis, while the other system involves a non-zero crossing angle in the x-axis. By using the aforementioned approach for both configuration systems, a system consisting of 22 equations can be derived.

By employing this methodology, it is feasible to achieve the necessary emittance using the LS, as illustrated in Figure (5.17b).

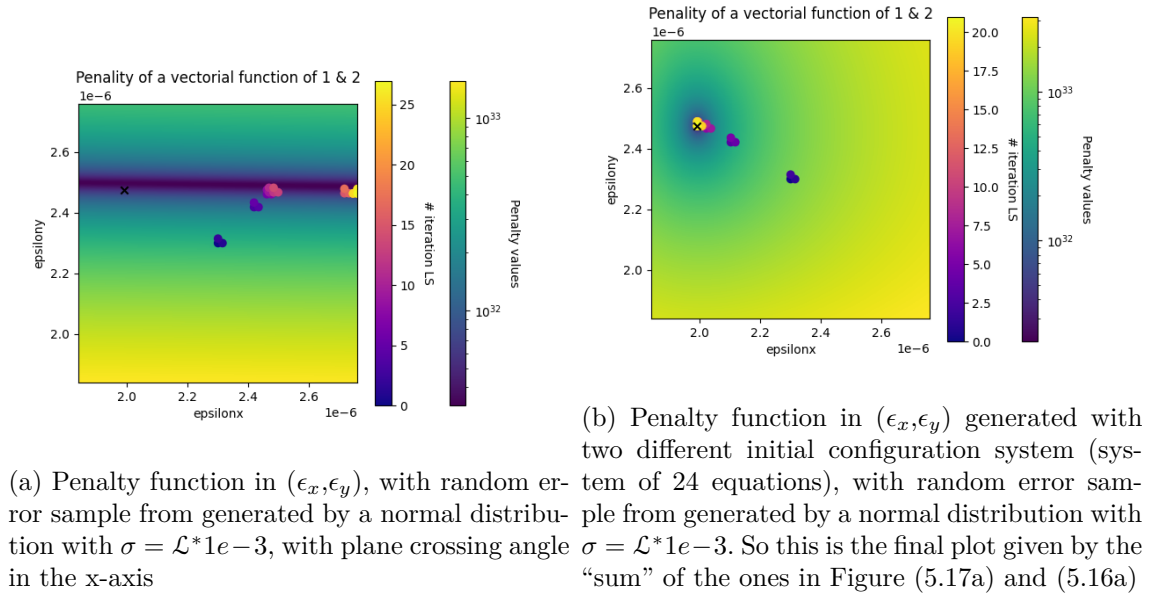


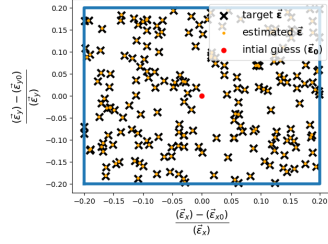
Figure 5.17: At left, the representation of the penalty functions with error of the order $1e-3$, with underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value, with plane crossing angle in the x-axis. At right the representation of the penalty functions with error of the order $1e-3$ from two initial configuration system with different plan of crossing angles, with underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value.

To ensure the validity of the methodology, the procedure was replicated for a total of two hundred distinct random selections of (ϵ_x, ϵ_y) . Figure (5.18) depicts the precision of this approach, considering shifting all the orbit parameters.

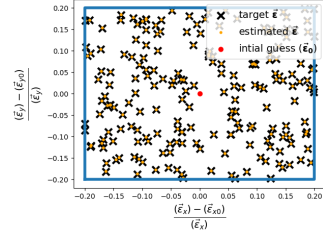
The Table (5.5) displays the average relative error and average calculation time for the least squares method across all 200 emittances, for all potential orbit shifts.

Different emittances for the two beams

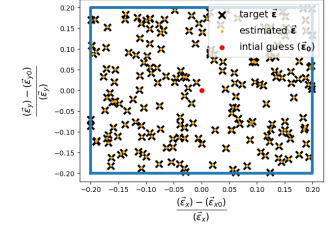
Unfortunately, as previously analysed also in this case ($\epsilon_1 = \epsilon_{x1} = \epsilon_{y1}, \epsilon_2 = \epsilon_{x2} = \epsilon_{y2}$) employing a single initial configuration and the described approach is inadequate to achieve the desired emittance. Exploiting that in the LHC there are different detectors, it is feasible to use two distinct initial configuration systems to achieve the desired objective. The challenge related to this technique is in identifying an alternative initial configuration that exhibits an identical shape to the one depicted in Figure (5.19a), but with a slightly curved “rift”.

Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the offset on y (μ_x)


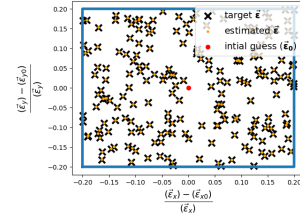
(a) Shift for the offset in the x-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the offset on y (μ_y)


(b) Shift for the offset in the y-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying crossing angle on x (θ_x)


(c) Shift for the crossing angle in the x-axis

 Reconstruction of target (ϵ_x, ϵ_y) by Luminosity, varying the crossing angle on y (θ_y)


(d) Shift for the crossing angle in the y-axis

 Figure 5.18: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_x, \epsilon_y)$ chosen randomly in the cyan square of 20% with respect to the initial guess, with different shifting parameter, considering a random error of 1e-3 on the luminosity value.

	(ϵ_x, ϵ_y) with err	LS time
μ_x	0.00010802083137094839	0.15717744827270508
μ_y	0.0001101968663368108	0.16458821296691895
θ_x	0.00010735539312379033	0.1674044132232666
θ_y	0.00011774767791462203	0.17041945457458496

Table 5.5: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, for all the possible shifts of the orbit.

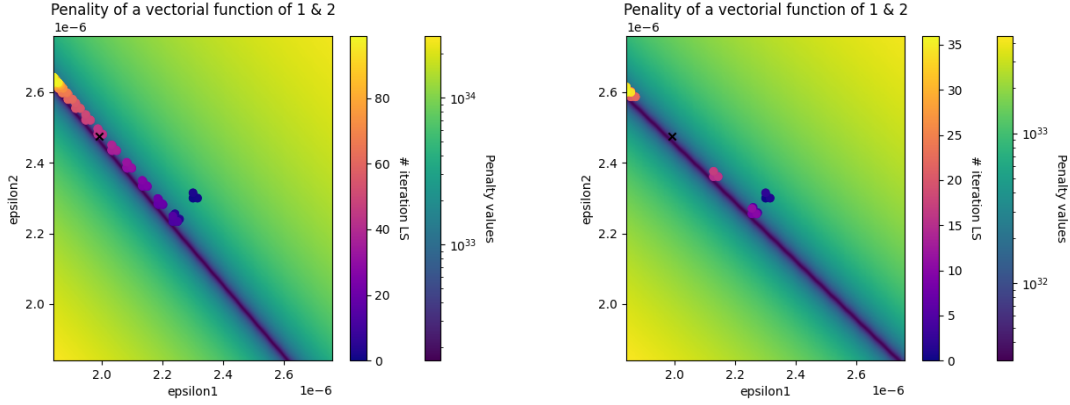
Based on the examination of the rift in section (5.2.2), it is intuitive to determine alternative configurations with varying proportions of β^* between the two beams. This intuitive approach is conclusive; indeed, in Figure (5.19b) it has been shown how it is modified the rift structure, changing the proportion of β^* between the two beams.

In the final penalty function, as depicted in Figure (5.20), the presence of the bore is less pronounced compared to Figure (5.17b). However, the steepness of the penalty function is sufficient to enable the least squares estimation to accurately determine the ideal emittance with a suitable level of error.

The approach has been evaluated by randomly selecting two hundred distinct emittances (ϵ_1, ϵ_2) values to assess its accuracy, as shown in Figure (5.21).

The Table (5.6) displays the average relative error and average calculation time for the least squares method across all 200 emittances, for all potential orbit shifts.

Unfortunately, the current approach is less achievable compared to the previous one. This is due to the fact that in the second configuration system, a difference of 20% has been introduced for the β^* values between the two beams. This difference exceeds the measurement error of this parameter.



(a) Penalty function in (ϵ_1, ϵ_2) , with random error sample from generated by a normal distribution with $\sigma = \mathcal{L}^* 1e - 3$, with the same β^* for both beams
 (b) Penalty function in (ϵ_1, ϵ_2) , with random error sample from generated by a normal distribution with $\sigma = \mathcal{L}^* 1e - 3$, with a difference for the β^* of the two beams of 20%

Figure 5.19: At left, penalty functions with the same β^* for both beams, underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value. At right, penalty functions with the different β^* for both beams (20%), underlined the ideal value of emittance and the different iteration of the LS in order to reach the goal value. Both plots of the penalty functions were with error of the order $1e-3$.

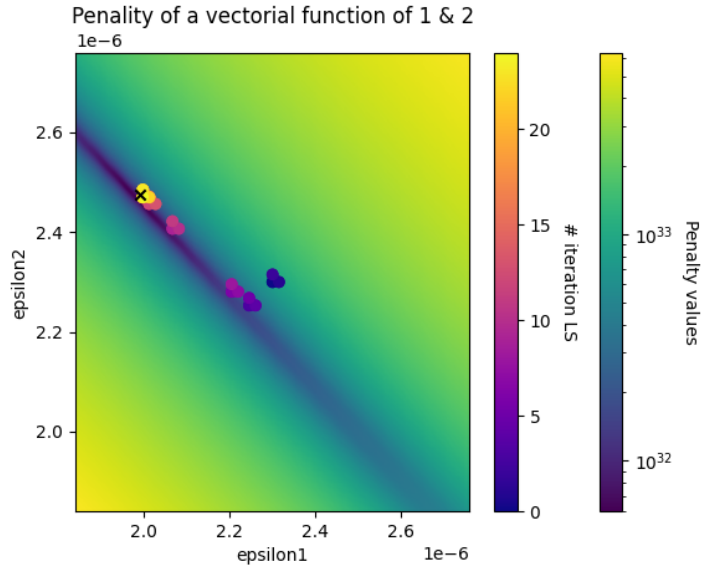
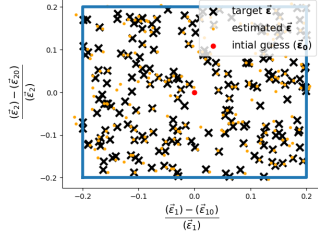


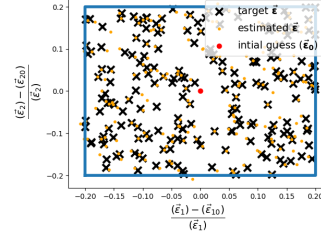
Figure 5.20: Penalty function in (ϵ_1, ϵ_2) generated with two different initial configuration system (system of 24 equations), with random error sample from generated by a normal distribution with $\sigma = \mathcal{L}^* 1e - 3$. So this is the final plot given by the “sum” of the ones in Figures (5.19a) and (5.19b)

5.3.4 4D case: realistic case

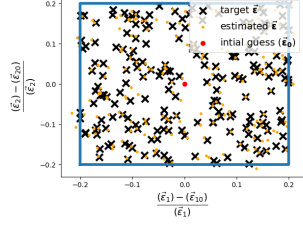
This particular case is the peak of realism and complexity examined within the context of this thesis. As in the beginning of the previous scenario, by employing the approach outlined

Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_y)


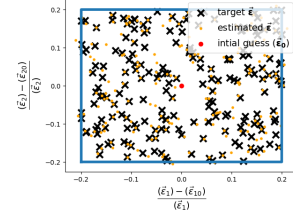
(a) Shift for the offset in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the offset on y (μ_y)


(b) Shift for the offset in the y-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying crossing angle on x (θ_x)


(c) Shift for the crossing angle in the x-axis

 Reconstruction of target (ϵ_1, ϵ_2) by Luminosity, varying the crossing angle on y (θ_y)


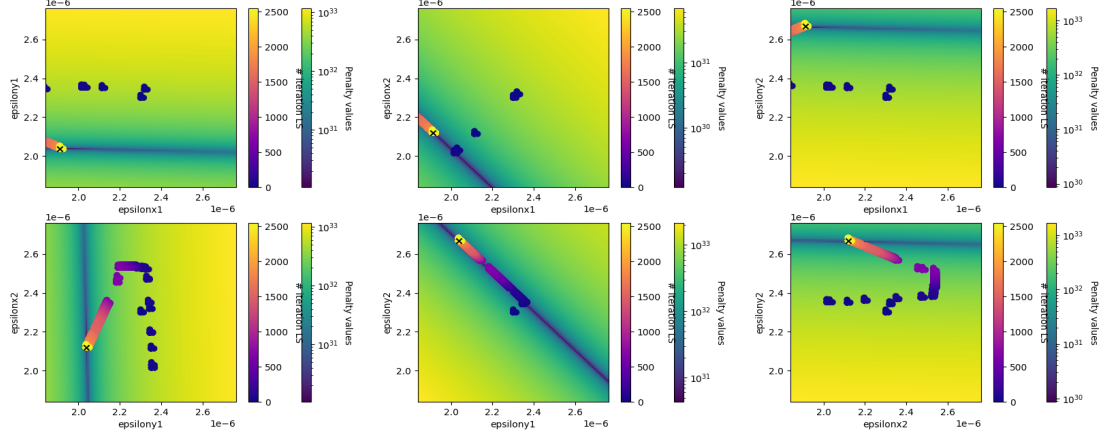
(d) Shift for the crossing angle in the y-axis

 Figure 5.21: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_1, \epsilon_2)$ chosen randomly in the cyan square of 20% with respect to the initial guess, with different shifting parameter, considering a random error of $1e-3$ on the luminosity value.

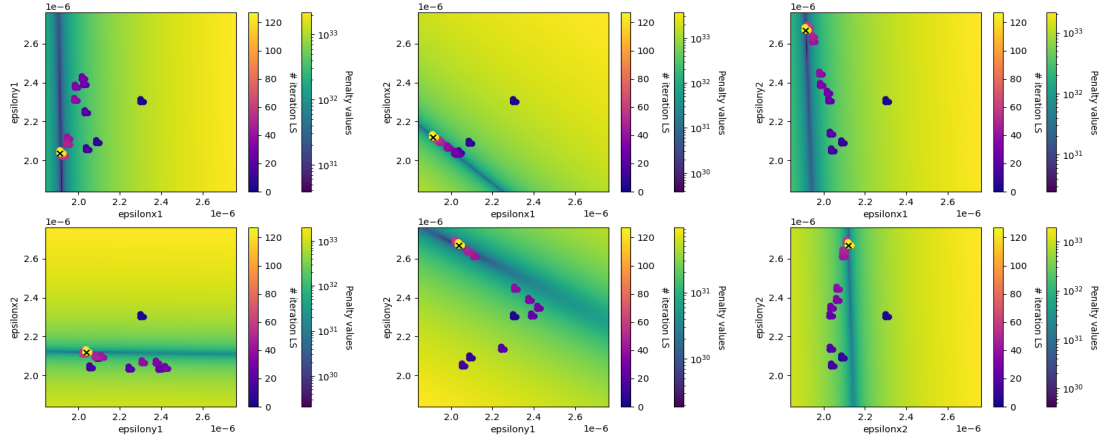
	(ϵ_1, ϵ_2) with err	LS time
μ_x	0.00517200003553423	0.17585420608520508
μ_y	0.005093828086478383	0.18398523330688477
θ_x	0.004551753308190394	0.18548011779785156
θ_y	0.004732640950103416	0.1875004768371582

Table 5.6: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, for all the possible shifts of the orbit.

and starting from a single initial configuration system, the LS algorithm successfully identified a minimum value for the penalty function, which deviated from the desired emittances. The idea of this case is to start the investigation by already considering two distinct configuration systems. The first configuration involves a plane crossing angle denoted as θ_x on the x-axis not null and an equal value of β^* for the two beams have. The second configuration exhibits a non-zero crossing angle on the y-axis (θ_y) and varying β^* values between the two beams, with a difference of 20%. As mentioned in the preceding section, without error, plotting the penalty function in this case presents a notable challenge due to its nature as a function in a five-dimensional space. Consequently, it has been determined that only the projections of the penalty function onto a two-dimensional plane, specifically around the initial guess of the LS, will be plotted. These projections will be visually represented using a colour code. The approach, starting from these two initial configuration systems, also inspiring on the projection plots of the error-free scenario originating from these two given configurations (as depicted in Figure (5.22)), appears to hold significant potential. This is due to the observed variations in the slope of the rift across all the different projections.



(a) Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$, with plane crossing angle on x-axis and equal β^* for both beams

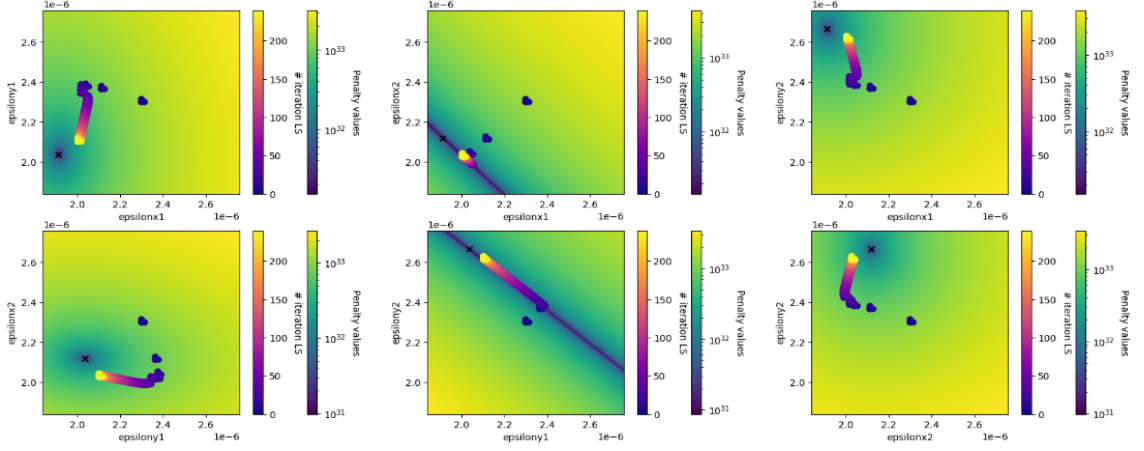


(b) Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$, with plane crossing angle on y-axis and different β^* between the beams, 20%

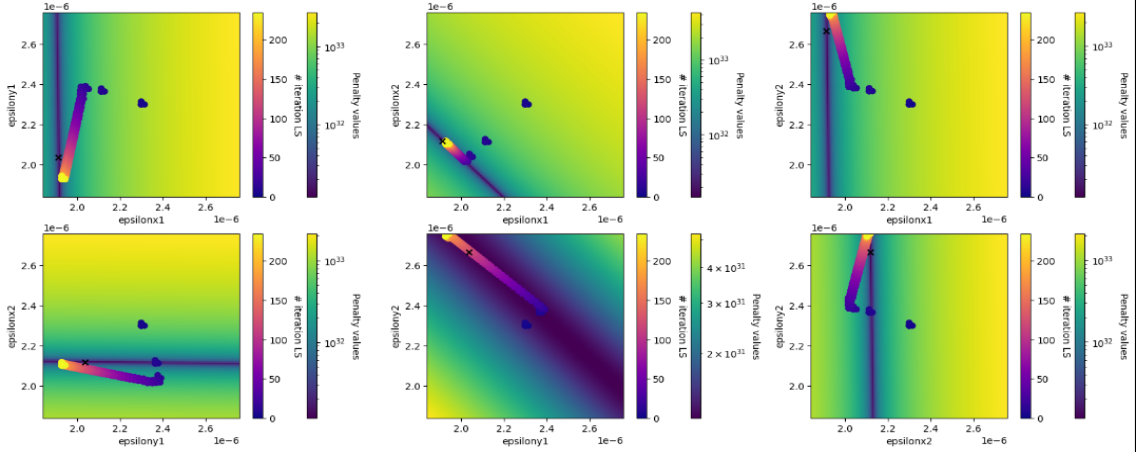
Figure 5.22: Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ in different initial configuration system.

Unfortunately, employing these two initial configuration systems and adopting the defined approach, which would result in a system consisting of 22 equations, proves to be insufficient, as depicted in Figure (5.23a). Based on the projections of the penalty functions for only the equations, with error, related to one of the two initial configuration systems, depicted in Figure (5.23b) and (5.23c), it can be concluded that the error in luminosity has a significant impact on their projections. Specifically, this effect is pronounced in the two epsilon values associated to the plane where the crossing angle deviates from zero in that configuration system. These observations align with the findings from the previous case analysis (5.3.3), where altering the slope of the rift was found to be a more challenging task. Consequently, the final “sum” of the two components in these two projection planes would result in a “sum” of two distinct rifts with differing widths. This would contribute to the overall penalty function, as depicted in Figure (5.23a), where the rift persists in these two

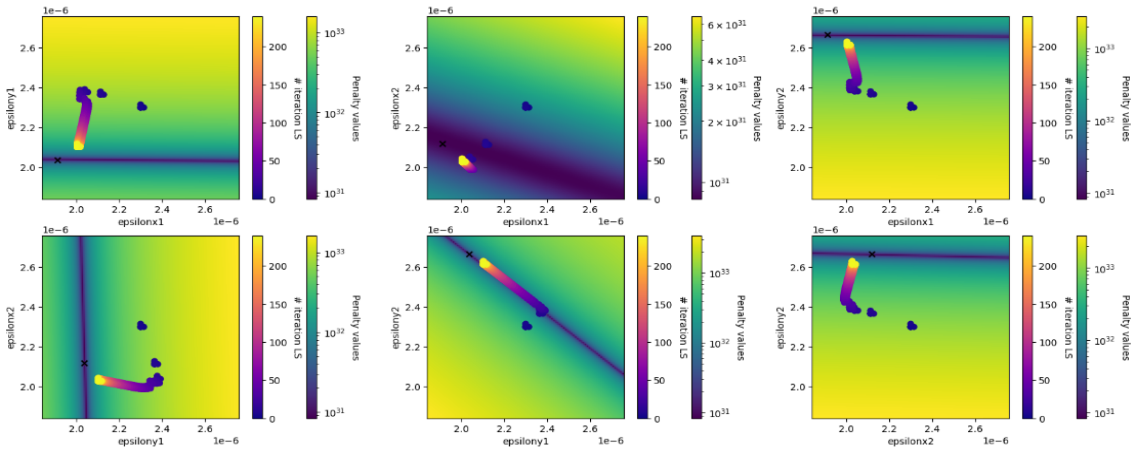
out of the six projections.



(a) Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ created by two different initial configuration system, (5.23b) and (5.23c)



(b) Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$, with plane crossing angle on y-axis and equal β^* for both the beams



(c) Penalty function for the different projections of the space $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$, with plane crossing angle on x-axis and different β^* between the beams, 20%

Figure 5.23: Projections of the penalty function on $(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ of a system, from different groups of initial configuration, with random error from a normal distribution with $\sigma = \mathcal{L}^* 1e - 3$.

As previously said, the purpose to have a penalty function that transforms a rift into a bore within the entire space is considered effective. In this case, the objective is to observe the presence of this bore in all the six projections of the five-dimensional plane. Nevertheless, this requirement is a necessary condition, but it may not be sufficient to determine if the LS achieves the desired emittance. This goal encounters similar challenges as attempting to analyse a two-dimensional plane by studying one-dimensional projections along the axis of a specific point. These projections, while necessary, are insufficient to comprehensively examine the entire plane.

One way to reach the objective is to distance even more from a realistic approach. The plots depicted in Figure (5.23) demonstrate that employing solely two initial configuration systems is insufficient for effectively “clarifying” the structure of the penalty functions and removing the occurrence of rifts leading to bores. From Figures (5.23b) and (5.23c), it is observed that there is a single projection, in each figure, which is particularly susceptible to errors. To address this issue, it is proposed to explore four distinct initial configurations. The first configuration involves a plane crossing angle (θ_x) on the x-axis not null and identical value of β^* for both beams. In the second configuration, the structure remains the same, but the plane crossing angle is now on the y-axis. The third configuration exhibits a non-zero crossing angle (θ_x) on the x-axis and varying β^* values between the two beams, differing by 20%. The last configuration shares the same structure as the previous one, with plane crossing angle on the y-axis. This choice would result in the formation of a solitary bore in each of the six distinct two-dimensional projections that have been previously depicted, as illustrated in Figure (5.24). It can be noticed that despite this condition, from a mathematical point of view, is not sufficient, the formation of these unique voids in the various projections, as a result of rotational symmetry, monotonicity, and other complex properties of the functions, is sufficient for the LS to attain the desired emittance, as depicted in Figure (5.24).

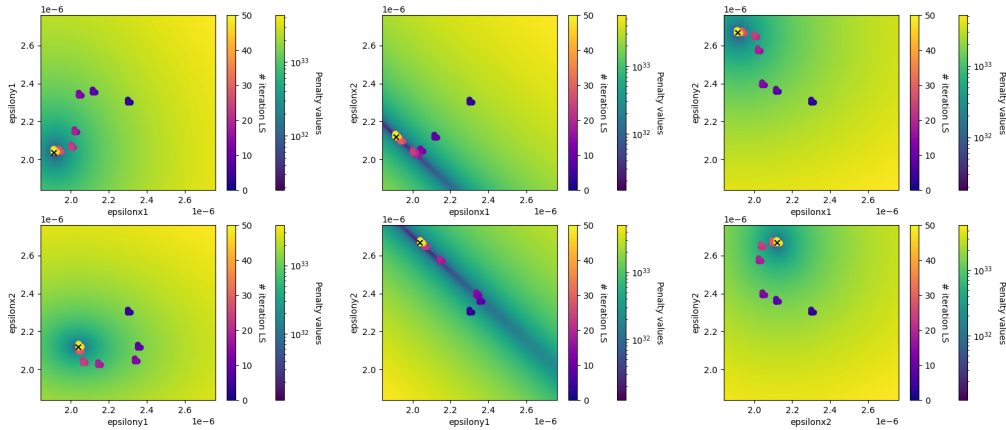


Figure 5.24: Penalty function for the different projections of the space ($\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2}$) created by four different initial configuration system, with random error sample from generated by a normal distribution with $\sigma = \mathcal{L} * 1e - 3$.

In order to improve the statistical validity of the analysis, a total of two hundred random examples of emittances were randomly selected and examined. The identification of these emittances was achieved just by employing shifts on the orbits, as depicted in Figure (5.25).

The Table (5.7) displays the average relative error and average computing time for the

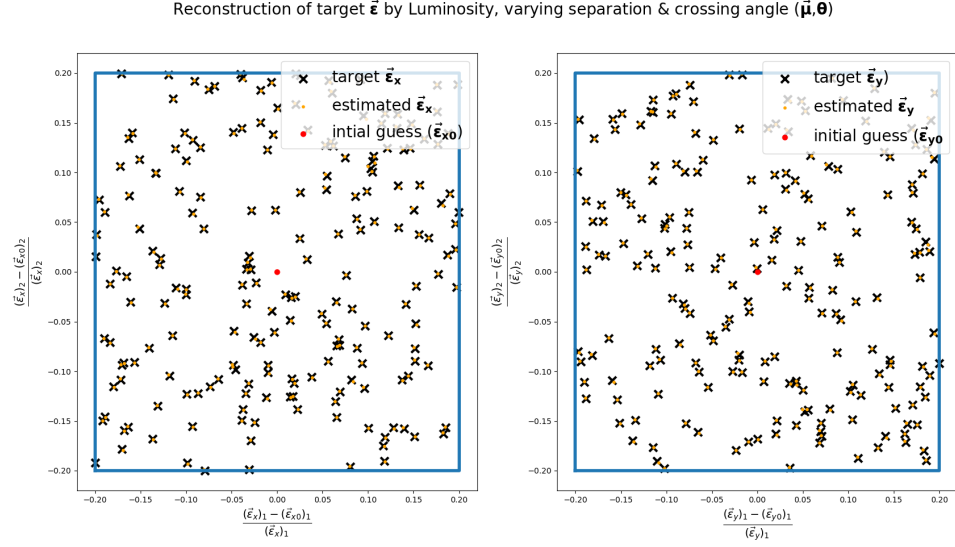


Figure 5.25: Accuracy of the approach for two hundred random emittances $\vec{\epsilon} = (\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ chosen randomly in the cyan square of 20% with respect to the initial guess, with different shifting parameter, considering a random error of 1e-3 on the luminosity value.

LS method across 200 different emittances, with all orbit parameters being shifted.

	$(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ with err	LS time
orbit	0.0011446379949476258	2.2011311054229736

Table 5.7: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, obtained shifting the different parameters of the orbit.

It is noteworthy to observe that the Table (5.6) indicates that the average error of the two-dimensional case ($\epsilon_1 = \epsilon_{x1} = \epsilon_{y1}, \epsilon_2 = \epsilon_{x2} = \epsilon_{y2}$) is larger compared to the four-dimensional case. This phenomenon may be attributed only to one factor, as it can be shown that the two problematic projections of the 4D case correspond to the planes studied in the 2D case mentioned. The solution for the 4D situation follows a similar approach to that of the 2D example, with the primary distinction being a greater number of equations. Specifically, the total number of equations in the system is doubled. This increase, for the law large numbers, serves to “reduce” the overall random error on the penalty function. Even a minor reduction of this error has a significant impact on accurately estimating the emittances value for the sensitivity of the approach.

Unfortunately, while this strategy may yield the desired outcome, it poses additional challenges in terms of implementing it in an actual accelerator. Specifically, already in the preceding situation (5.3.3), it was challenging from a physics perspective to achieve significantly different β^* values for both beams within the same detector. The current objective has become even more intricate. It necessitates either the use of four distinct detectors, each starting with a unique initial configuration system, or the use of only two detectors with the ability to switch the crossing angle during operation while maintaining a constant β^* , or alternatively modifying the β^* value for one of the two beams while keeping the plane crossing constant. It is evident that if the preceding scenario was challenging to implement,

the current one is far more laborious.

One notable aspect of this approach is its composition of multiple equations and its small error, as demonstrated in Table (5.7). As a result, some assumptions regarding the starting configurations can be relaxed. It is always essential to have four initial configurations. However, if we examine a favourable scenario where the change in β^* between configurations is 2%, the average relative error of 200 random samples is presented in the Table (5.8).

	$(\epsilon_{x1}, \epsilon_{y1}, \epsilon_{x2}, \epsilon_{y2})$ with err	LS time
orbit	0.006575467531951761	2.207535743713379

Table 5.8: Table representing the average for all the 200 samples emittances of the relative error and computation time of the non liner LS estimation, obtained shifting the different parameters of the orbit

When the change in β^* is further relaxed by 2‰, the resulting error is approximately 7.5%.

5.4 Conclusions and following studies

This chapter discusses the utilisation of the luminosity model and its optimized implementation, as previously introduced in Chapter (4.2), to address an ill-posed problem. Specifically, the objective is to determine the transverse emittances based on a singular scalar, namely the luminosity.

This chapter provides a comprehensive overview of various methodologies and approaches to address the challenges encountered during the analysis process. It specifically focuses on handling complex situations and highlights the favourable outcomes achieved in different scenarios, in contrast to the existing instrument employed for emittance measurement.

Various instances have been evaluated in an attempt to identify the most realistic one, that takes into account the nominal random error associated with the luminosity value. In each of the situations discussed, a strategy was identified that enabled the attainment of the reasonable level of error in the estimation of emittance.

Undoubtedly, the aforementioned solutions exhibit varying degrees of feasibility depending on the specific circumstances. It would be of great interest to explore other strategies that could potentially yield comparable or superior outcomes to those presented, while avoiding the limitations inherent in our current methodologies. The latter two cases mentioned would be the most intriguing to improve.

Considering that the idea behind the approaches employed in this chapter aligns with the concept presented in the introduction of Section (5.2).

It may be of interest to remove another hypothesis that is assumed in Section (5.2), the one pertaining to the precise luminosity model. In relation to this, it is well-established that the current model used is not yet complete. Specifically, it is important to consider in the model the inclusion of the crab cavities' contribution ([Ca12]), as well as the impacts stemming from beam-beam interactions ([He14]) and so on. One potential enhancement could involve incorporating the model of these phenomena into the luminosity formula, such as incorporating generic errors that are dependent on the emittance, the latter improvement would allow for the consideration of particularly intricate effects that are challenging to model. By doing so, the approach, described in (5.2), of examining the collision from several perspectives is expected to still yield positive results. However, it is possible that alternative strategies may be required to deal with the ill-conditioning of the problem. It would

be intriguing to evaluate the impact of the least squares method in terms of both relative error and computational time for this specific computational problem.

Another potential improvement is removing the remaining hypothesis pertaining to the accurate measurement of the parameters, as discussed in Section (5.2). In this particular scenario, the values of the parameters can be affected by two types of errors:

- Random error: in this particular scenario, the issue at hand is similar to the previously examined problem. The presence of a random error in the parameter has the potential to impact the value of luminosity, which in turn may be subject to a related random error. Upon initial examination, it appears that the approach described in (5.3) may be applicable also to this particular scenario. However, conducting a thorough investigation would be highly intriguing in order to validate this assumption.
- Systematic error: in this case, the problem is completely different. Now the error is due to calibration problems, so now the error on the selected parameter doesn't change in each equation, but changes its impact on the luminosity value. One potential solution is treating the error as an unknown variable within the LS method, which allows a direct estimation of the error from the calculation process. Unfortunately, the LS system, which encompasses the manipulation of all orbit parameters, consists of five distinct equations. Therefore, it would be capable of estimating a systematic error for a single parameter within the configuration system. It would be intriguing to investigate whether employing the same approach as described in (5.3) would result in an improvement of the least squares estimation when systematic errors are present in several parameters. If not, it is intriguing to examine this scenario as well, in order to comprehensively understand the problem in its entirety.

Lastly, it would be intriguing to explore the aforementioned analysis by releasing the assumptions regarding Gaussian beam profiles (e.g. considering q-Gaussian ones, used in ([Pa20])) and the factorization of the beam profile distribution for the x-y axes.

Conclusions

As introduced in Section 1.2, the luminosity for each given Interaction Point can be factorized by calculating the cumulative luminosity resulting from all the collisions that occurred at that specific location

$$\mathcal{L}_{IP} = n_{\{bb,IP\}} \mathcal{L}_{bb}.$$

The objective of this thesis was to optimise the luminosity in a comprehensive manner by effectively using the aforementioned factorization. The initial two chapters were dedicated to analysing the number of bunches, $n_{\{bb,IP\}}$, colliding in a specific detector, while the remaining two chapters focused on the instantaneous luminosity bunch-by-bunch, $\mathcal{L}_{bb}$. The fundamental idea of this two-step investigation presented similarities in the approach adopted for each step:

- The odd chapters of the study mostly centred on the physics and the underlying model (analytic phase). Additionally, these chapters examined the potential for optimising a pre-existing Python implementation used at CERN by employing specific techniques to reduce its computational costs. In particular, it has been used for:
 - the number of collisions: an alternative model has been employed, resulting in an improvement in performance by ~ 100 compared to the previously used model
 - the luminosity bunch by bunch: it has been employed using the same model, however with different implementations aimed at boosting its performance through the utilisation of Numba. This optimisation has resulted in a performance improvement of up to a factor of 100.
- The primary objective of the even chapters was to investigate the optimization of the inversion problem by a synthetic approach, supported by the given model:

- In the second chapter, the studies focused on the impact of the filling scheme on the number of collisions. Specifically, the goal was to optimise the arrangement of bunches in order to maximise the number of collisions.

The chapter presented an interesting outcome in the form of a parallel implementation of a Monte Carlo simulation. This simulation was specifically tested on the widely employed filling scheme in the LHC in 2022. The results of the simulation revealed several feasible alternative filling schemes with a larger number of collisions compared to the scheme employed in the accelerator (up to 2% gain). Moreover, this simulation has been used either for ions operational phase in 2023 and in some experimental session, namely Machine Development, with protons in 2023. In this study, we assume filling schemes composed exclusively of a single type of PS batch. To generalize the present results, one could consider

- * to apply the same analysis for hybrid filling schemes: filling scheme composed by different PS batches

- * to extend the Monte Carlo simulation: not only shifting the SPS batch, but also the PS batch
- * to improve the parallelization of the Monte Carlo simulation
- In the fourth chapter, the bunch by bunch luminosity (scalar) has been inverted, using a perturbation theory approach, to derive some variables of the model, specifically the four transverse emittances of the colliding beams. This problem is clearly ill-conditioned, because the goal is to obtain a vector from a scalar. In the chapter it has been introduced some strategies and approach in order to condition the problem, also exploiting the python optimization obtained from Numba exposed in the third chapter. The analysis of this scenario also took into consideration the existence of noise in the luminosity determination. The primary objective of this chapter was to describe various ways for successfully acquiring the emittances from the luminosity, and, at the same time, to outline the limitation associated with these particular approaches.

From the present status of the study, we see several interesting studies worthy of exploration:

- * Identify alternative numerical methodologies that yield comparable outcomes to the one elucidated, while overcoming the limitation delineated in the chapter.
- * To properly consider that the model used is not fully accurate to describe the phenomenon. In particular, that are some known effects that are not considered in the model like the beam-beam or the Crab Cavities, it would be interesting to see how change the performance of the approaches presented in this generalization.
- * Consider that the values of the different parameters measured from the machine may be subject to both random and systematic errors.
- * Relax the hypotheses on the Gaussian profiles (e.g., considering q-Gaussian ones) and on the x-y factorization of the bunch transverse distributions.

Appendix A

A.1 Studies on Long-Range interactions

In this appendix is shown an approach to obtain a filling scheme of the two beams that would lead to a filling pattern with Long Range interaction only in some specific slots near an interaction point

A.1.1 Using the mathematical model...

The purpose of this theorem is to create a pattern that has to be respected, from a mathematical point of view, in order to get Long Range collisions only in some specific slots.

Theorem A.1.1. *Given two beams, the first only with N consecutive bunches and the second one with M bunches, which positions are set in order that their first particle is symmetric with respect to the slot where all the collisions happen. There will be $N-M+1$ points where it happens at least one collision in the semicircle where all the N bunches stood before the collision (either Head On and Long Range).*

For notation we will call these points as interaction slots, because considering that the two beams are moving one against the other, these interaction slots are spaced by the half of the accelerator slot, so in the LHC they are spaced by ≈ 3.75 m.

Proof. In the proof, that will be per induction, the choices of 0 as position of collisions and to fix N bunches on left of 0, are done for sake of simplicity, but the position of the collisions can be shifted, and we could fix the particles on right of that position.

1. So, for simplicity, we fix the number of bunches on left of 0, N_b on left (then all empty slots) and 1 beam on right of 0 for the two beams, respectively. From construction, it's trivial that there will be N ($N-1+1$) interaction slots where it happens at least 1 collision, the one Head On in slot 0 all the other Long Range on left of slot 0 (for notation this N interaction slots where it happens at least 1 collision we will call them as true interaction slots and the ones where didn't happen false interaction slots). In order to get easier the following steps, now we consider the case of N bunches on left of 0 and 2 on right. The 2° bunch of the 2° beam will collide with all the N bunches of the other beam, but with one time in delay with respect to the 1° bunch of its own beam (12,5 ns later). Then, considering that we had a vector for the collisions with the 1° bunch of the 2° beam, the collisions with the 2° bunch is the same vector but shifted by one interaction slot. So in this case we will have $N+1$ collisions: the Head on at 0, and for the Long range $N-1$ on left of 0 and one on right.
2. If we had N bunches for the 1° beam and M for the 2° beam, and I suppose to have $N-M+1$ true interaction slots in total, now if I add another bunch at the end of the queue of the M bunches for the 2° beam I will have that this last will collide with N bunches of the 1° beam but with one interaction slot shifted with respect to the vector of collisions for M° bunch of the second beam, so with respect to the case N_b and M_b now we will have $N-M$ true interaction slots.

□

A.1.2 ... in a practical point of view

Considering that in the LHC (in proton studies, in order to maximize the collisions in ATLAS/CMS) the disposition of the two beams will be the same and symmetric with respect to the origin, we can model every kind of bunches disposition used in the LHC as a

NbMe, that is composed by N consecutive bunches (spaced every 25 ns) and M consecutive empty spaces (accelerator slots), continuing to repeat always the same structure. Then, considering the two beams modeled as a NbMe on left and right of the origin, we know from the theorem that the collisions of the bunches between the two beams are the same as before but now the bunches will meet also the empty spaces. For sake of simplicity, we will see only the collisions of the bunches of beam 1 but it can be extended for symmetry also in the other one. Considering the algebraic structure behind the theorem, it can be extended easily for the match of bunches and empty spaces, so as before the first that runs into all the M empty spaces will be the 1° bunch of the beam 1. Then, we know that from the $(N)^\circ$ interaction slot there will be no collisions **between bunches** for M consecutive interaction slots (considering also the N° interaction slot); instead the 2° of beam 1 will have the same collisions but with one interaction slot of delay, but in this case we cannot see this delay because in the $(N-1)^\circ$ interaction slot is already true, there has already been a collision between the 1° bunch of beam 1 with the $(N)^\circ$ bunch of beam 2, so nothing changed in the output of interaction slots. Then we iterate this structure for all the matches between the bunches of beam 1 with the empty spaces of beam 2. As result we will have false interaction slots in these intervals $[-M-N+1, -N+1]$, $[N-1, M+N-1]$ and true ones in $[-N+1, N-1]$. At the end, considering the two beams as NbMe repeated consequently until the end of the space, after the matches between Nb of beam 1 and Me of beam 2, there will be another collision with Nb between the two beams in the $(M-N)^\circ$ slot, that will follow the theorem because the symmetric hypothesis is respected, so we will have $N-1$ true interaction slots on right and left of $(M-N)^\circ$ slot. So eventually we will have, considering all symmetric on left of 0, true interaction slots in the intervals $[0, N-1]$ and $[M+1, M+2N-1]$, instead of false interaction slots between $[N, M]$ and then this pattern repeats continuously until all the space available is filled, and all this structure is repeated around IP5 for symmetry.

A.2 Optimized algorithm Pre-fill without INDIV

Here it is represented the different steps, for the computation of the "best" filling scheme in the most general case without INDIVs.

1. The following notation will be used. Let b be the quantity of bunches for each PS batch. The typical number of vacant spots between successive PS batches (*empties*) is 8. The variable SPS_{batch} represents the maximum length of an SPS batch, which is defined as the number of PS batches included within it. Additionally, the variable inj_{space} denotes the amount of empty space between SPS batches, often set to a value of 32. The variable bun_{tune} represents the quantity of bunches utilised for the tune measurements, typically set at a value of 12. inj_{tune} is the slot number where to inject the other beam for the tune in order to prevent collisions between the two beams in ATLAS and CMS, which is typically 12, given that in one beam this slot is 0.
2. The initial step in this process, drawing inspiration from the algorithm discussed in the preceding chapter, involves determining the optimal length of SPS batches. This is necessary to maximise the number of bunches within both rings:
 - In this study, we introduce the following notation: n_i represents the length of the SPS (Super Proton Synchrotron) that we wish to test, where $1 \leq n_i \leq n$. Here, n denotes the maximum SPS length specified by the user. Additionally, we define M_{ni} as the number of SPS batches with length n_i in a given quarter apart the first one, and p_{ni} as the length of the last SPS in that quarter. Furthermore,

M_{ni1} represents the number of SPS batches with length n_i in the first quarter, while p_{ni1} denotes the length of the last SPS in the first quarter. The variable bun represents the number of bunches in each PS (Proton Synchrotron) batch, as determined by the user. Finally, we introduce q as the number of free slots in the first quarter between the bun_{tune} and the first train.

- from the definition it's possible to compute:
 - $M_{ni} = \text{int}(\frac{891}{bun(n_i)+7(n_i-1)+31})$ with $M_{ni} \leq M_{ni-1,1} \forall n_i > 1$
 - $M_{ni1} = \text{int}(\frac{891-121-inj_{tune}-bun_{tune}-9}{bun(n_i)+7(n_i-1)+31})$ and
 $p_{ni1} = \text{int}(\frac{891-(31+bun)-121-(inj_{tune}+bun_{tune})-9-(bun(n_i)+7(n_i-1)+31)M_{ni1}}{bun+7})$, with $M_{ni1} \leq M_{ni-1,1} \forall n_i > 1$, $p_{n1} \leq n-1$ and $p_{ni1} \leq n \forall n_i < n$
 - $q = \text{int}(\frac{891-121-inj_{tune}-bun_{tune}-9-M_{ni1}(bun(n_i)+7(n_i-1)+31)-p_{ni1}*n_i-(p_{ni1}-1)7}{3})$ and
 finally $p_{ni} = \text{int}(\frac{891+3+q-(bun(n_i)+7(n_i-1)+31)M_{ni}}{bun+7})$, with $p_n \leq n-1$ and $p_{ni} \leq n \forall n_i < n$
- the total number of bunches in the quarter are: $bun[n_i M_{ni} + p_{ni} + 1]$

The algorithm yields a vector with the varying number of bunches within the quarter for all possible choices of ni . The greatest value can be selected;

3. Given the current constraint regarding the abort gap, it is necessary to adjust the analysis of the quarters by accounting for the length of the abort gap, which consists of 121 slots. Consequently, the first quarter will be smaller than the subsequent ones due to the presence of the abort gap, which prohibits injection. It is important to note that the structure under consideration is periodic in nature, as it forms a ring. In the absence of explicit notation, we shall consistently refer to the divisions as "quarters." However, it is important to note that these quarters are not precisely aligned with the quarter of the ring, with ATLAS serving as the reference point (designated as 0). Each quarter is comprised of a total of (891-121,891,891,891) slots.

Hence, the algorithm seeks to use the altered structure of these quarters by populating the Abort Gap Keeper with the most optimal train calculated previously. It then proceeds to fill all subsequent quarters with the same train, excluding the last one (which has a shorter length denoted as p_{ni}). It is also feasible to insert the bun_{tune} at the start of the ring, namely in slot 0, for one of the two beams. Similarly, the other beam can be placed in the slot determined by the inj_{tune} .

4. Thirdly, the structure of the previous quarter is replicated and applied to the remaining quarters. Each quarter is shifted by three slots, resulting in a cumulative shift of nine slots for the first quarter. This adjustment is made to align the quarters harmonically with the Large Hadron Collider (LHC). Specifically, the position of the LHCb detector in the convolution was 2670 on the "left" for Beam 1. Consequently, it needed to be shifted by 894 slots to the "right" ($3564-2670 = 894 = 891+3$) to match the desired harmonic position. This copy and paste procedure was necessary to optimise the utilisation of SPS batches within this harmonic configuration.
5. The shift has caused substantial modification to the harmonic frequencies. Nevertheless, it is important to acknowledge that a significant proportion of bunches still correspond to the desired harmonics of the target detectors. The primary goal of this stage is to strategically use the available vacant time slots in order to include additional shorter trains in an optimized manner. The methodology employed to accomplish this objective is outlined as follows:

- After duplicating the most superior trains across the rings, we have the ability to insert the smaller train in the initial quarter, situated between the clusters of melodies and the foremost exemplary train. Despite the administration of the injection, it is possible that residual vacant space remains. This is due to the fixed number of bunches in the shorter train, which may result in unoccupied slots inside the initial "quarter" of the train.
- The subsequent procedure for utilising the available space involves evenly distributing it across one quadrant of the ring. How might this be accomplished? The length is divided by three and the resulting quotient is used to shift the different quarters. This process is performed in order to distribute the available space evenly throughout the quarters, effectively filling these unoccupied slots in the first quarter while creating three distinct free areas between the quarters. The combined lengths of these free areas is equal to the length of the initial free space. The rationale behind this approach is that by adjusting the positioning of the quarters, the LHCb harmonic is disrupted, resulting in a reduction in the number of collisions. However, by increasing the spacing between the quarters, there is a potential opportunity to introduce an additional PS batch within the shorter SPS batch (from $p_{ni} + 1$ to $p_{ni} + 2$), thereby increasing the number of collisions in ALTAS/CMS and maybe also again in LHCb.

If the inclusion of this passage does not result in an increase in the number of PS batches in the shorter train, we can proceed to reposition the quarters as described in step 4. By doing so, it becomes feasible to inject the shorter train, which has a length of $p_{ni} + 1$, between the quarters in the available space. This injection is carried out in a manner that ensures the slots where the SPS batches are inserted between consecutive quarters are separated by a distance of 894 slots ($891+3$), in order to respect again the LHCb harmonic.

The algorithm can be readily expanded to optimize ALICE, with the key distinction between ALICE and LHCb being the placement. Specifically, only one step in the algorithm needs to be modified, which pertains to the harmonic positions of the SPS batches. Given that ALICE's position with respect to the 4th harmonic obviates the necessity for a continuous shift of 3 slots, this step is replaced with a conventional copy and paste procedure that no longer involves shifting.

The algorithm shown above is a logical progression of the previously derived mathematical solution, incorporating certain modifications to account for the influence and limitations imposed by the additional restrictions introduced in the optimization issue.

Bibliography

- [CERN] Keijo, M.: “*CERN Council WebSite*”. CERN Council WebSite.
- [AT08] The ATLAS Collaboration: “*The ATLAS Experiment at the CERN Large Hadron Collider*”. JINST, (2008)
- [CM08] The CMS Collaboration: “*The CMS Experiment at the CERN LHC*”. JINST, (2008).
- [Aad12] Aad,G.: “*Observation of a New Particle in the Search for the Standard Model Higgs Boson with the ATLAS Detector at the LHC*”, Physics Letters B, pp. 1–29, (2012) <https://doi.org/10.1016/j.physletb.2012.08.020>
- [AL08] The ALICE Collaboration: “*The ALICE Experiment at the CERN LHC*”. JINST, (2008).
- [LH08] The LHCb Collaboration: “*The LHCb Experiment at the CERN LHC*”. JINST, (2008).
- [We16] Wenninger, J.: “*Machine protection and Operation for LHC*”. CERN Yellow Reports, p. 377, (2016). <https://doi.org/10.5170/CERN-2016-002.377>
- [Hal11] Halkiadakis, E.: “*Introduction to the LHC Experiments*”. TPhysics of the Large and the Small, pp. 489–518, (2011) <https://doi.org/10.1142/97898143271830009>
- [Wi17] Winn, M.: “*Prospects for quarkonium measurements in p - A and A - A collisions at the LHC*”. Few-Body Syst., (2017)
- [Da18] Damerau, H., Hancock, S., Lasheen, A., Perrelet, D.: “*RF manipulations for special LHC-type beams in the CERN PS*”. Proceedings of the 9th International Particle Accelerator Conference, Vancouver, Canada, (2018) <https://doi.org/10.18429/JACOW-IPAC2018-WEPAF063>.
- [Ga01] Garoby, R.: “*Multiple bunch-splitting in the PS: Results and plans*”. 11th Workshop on LEP Performance. (2001).
- [Bo99] Boussard, Daniel and Linnecar, Trevor Paul R: “The LHC Superconducting RF System”. LHC-Project-Report-316, CERN-LHC-Project-Report-316(1999).
- [Ve17] Velotti, F M and Barnes, M and Bartmann, W and Bartosik, H and Bracco, C and Carlier, E and Damerau, H and Fraser, M A and Goddard, B and Kain, V and Kotzian, G and Kramer, T and Schwick, C and Trad, G: “SPS and LHC Batch Spacing

- Optimisation"8th Evian Workshop on LHC beam operation, 137-140(2017).
- [Jo18] Jones, Rhodri: "Measuring Tune, Chromaticity and Coupling"2018 CERN–Accelerator–School course on Beam Instrumentation, (2018).
- [We17] "Wiesner, C and Bartmann, W and Bracco, C and Carlier, E and Fraser, M and Goddard, B and Jacquet, D and Magnin, N and Wenninger, J: "Abort Gap Keeper and Abort Gap Protection"8th Evian Workshop on LHC beam operation,(2017).
- [Vidal] Vidal, X.C., Manzano, R.C.: "*Buckets and Bunches*". Taking a closer look at LHC. <https://www.lhc-closer.es/takingacloserlookatlhc/0.bucketsandbunches>.
- [Hos18] Hostettler, M., Fuchsberger, K., Papotti, G., Pieloni, T., Papaphilippou, Y.: "*BLuminosity Scans for Beam Diagnostics*". Physical Review Accelerators and Beams, vol. 21 (2018). <https://doi.org/10.1103/PhysRevAccelBeams.21.102801>.
- [LPC] LHC Programme Coordinator. <https://lpc.web.cern.ch/schemeEditor.html>
- [La16] Lamont, M.: "*LHC Report: playing with angles*" CERN Accelerating Science. <https://home.cern/news/news/accelerators/lhc-report-playing-angles>
- [He01] Herr, W., Zorzano, M.P.: "Coherent dipole modes for multiple interaction regions". LHC Project Report 462 (2001).
- [Jo99] Jowett, J. M.: "Filling schemes, collision schedules, and beam-beam equivalence classes". SIS-2001-088 (1999).
- [FillPatt] Iadarola, G., Poyet, A., Sterbini, G.: "*Tool for the analysis of LHC filling schemes*" <https://github.com/PyCOMPLETE/FillingPatterns>
- [CollSched] Rufolo, M.: "*Tool for the analysis of LHC filling schemes*" <https://gitlab.cern.ch/mrufolo/fillingstudies/-/blob/master/README.md>
- [npbits] Numpy packbits <https://numpy.org/doc/stable/reference/generated/numpy.packbits.html>
- [Co16] Coupard, J and Damerau, H and Funken, A and Garoby, R and Gilardoni, S and Goddard, B and Hanke, K and Manglunki, D and Meddahi, M and Rumolo, G and Scrivens, R and Chapochnikova, ELHC Injectors Upgrade, Technical Design Report, CERN, (2016)
- [Fa21] Fartoukh, S and Kostoglou, S and Solfaroli C, M and Arduini, G and Bartosik, H and Bracco, C and Brodzinski, K and Bruce, R and Buffat, X and Calviani, M and Cerutti, F and Efthymiopoulos, I and Goddard, B and Iadarola, G and Karastathis, N and Lechner, A and Metral, E and Mounet, N and Nuiy, F X and Papadopoulou, P S and Papaphilippou, Y and Petersen, B and Persson, T H B and Redaelli, S and Rumolo, G and Salvant, B and Sterbini, G and Timko, H and Tomas G, R and Wenninger, J LHC Configuration and Operational Scenario for Run 3, CERN, (2021)
- [Zi02] Zimmermann, F Electron cloud effects in the LHC, CERN, (2002)

- [Se92] Sessler A.M.: Accelerator beam emittance, American Journal of Physics, (1992)
- [Hi21] "Hillert, W.,: 'Transverse Linear Beam Dynamics', arXiv, (2021).
- [Bu94] "Buon, J.,: 'Beam phase space and emittance', CERN, (1994).
- [St20] "Sterbini, G.,: 'Linear Optics Computations', arXiv, (2020).
- [dgaranin] "Garanin, D.,: 'Hamiltonian Mechanics' HamiltonianMechanics.pdf.
- [So88] Sørensen, A H: "Liouville's theorem and emittance", CERN–Accelerator–School course: General Accelerator Physics, CERN, (1988).
- [Twiss] Phase Space Distributions and Emittance in 2D Charged Particle Beams
<https://www.comsol.fr/blogs/phase-space-distributions-and-emittance-in-2d-charged-particle-beams/>
- [Ku16] Kuhn, M: "Transverse Emittance Measurement and Preservation at the LHC" CERN, (2016).
- [Ca09] Caspers, F: "Schottky signals for longitudinal and transverse bunched-beam diagnostics" CERN, (2009).
- [Mi03] M. Minty and F. Zimmermann, Measurement and Control of Charged Particle Beams, 99-130, Springer, 2003.
- [Ol13] Olson T., Fraser, M A and Lanaia, D and Voulot, D.: "Three-gradient Emittance Measurements at REX-ISOLDE", HIE-ISOLDE-PROJECT, (2013).
- [St15] J. W. Storey et al.: "Development of an Ionization Profile Monitor Based on a Pixel Detector for the CERN Proton Synchrotron", Proceedings of IBIC2015, (2015).
- [Mu06] "Muratori, B and Herr, W.,: "Concept of luminosity" CERN, (2006).
- [Co14] "Coombs, G.R.: "Luminosity and Luminous Region Calculations for Different LHC Levelling Scenarios" Imperial College, (2014).
- [Br91] "Brown, K. L. and Servranckx, R. V.: "Luminosity and Luminous Region Calculations for Different LHC Levelling Scenarios" Particle Accelerators, Vol. 36, pp.121-139, (1991).
- [Sy07] M. Syphers, "Some notes on Luminosity Calculations", Fermilab note, Beams-doc-1348 (2007).
- [Pa20] S. Papadopoulou and F. Antoniou and T. Argyropoulos and M. Hostettler and Y. Papaphilippou and G. Trad, "Impact of non-Gaussian beam profiles in the performance of hadron colliders", Physical Review Accelerators and Beams, (2020).
- [Ba21] Balagura V, "Van der Meer scan luminosity measurement and beam-beam correction", The European Physical Journal C, (2021).

- [Wh10] White, S and Alemany-Fernandez, R and Burkhardt, H and Lamont, M, "First Luminosity Scans in the LHC", CERN, (2010).
- [LS] Documentaion of Non Linear Least Squares in scipy <https://scipyLS>
- [Numba] Numba documentation <https://numba.pydata.org/numba-doc/latest/index.html>
- [PythonLumi] Documentation of python implementation of Luminosity computation <https://luminosityformula>
- [Ca12] Calaga, R, "Crab Cavities for the LHC Upgrade", CERN, (2012).
- [He14] Herr, W. and Pieloni, T., "Beam-Beam Effects", CERN, (2014).