

École polytechnique de Louvain

Exploration of dose trade-offs for proton therapy treatments using deep learning models

Author: **Marine VAN RENTERGHEM**

Supervisors: **Ana María BARRAGÁN MONTERO, John LEE, Jimmy OLSSON**

Readers: **Christophe DE VLEESCHOUWER, Margerie HUET**

Academic year 2021–2022

Master [120] in Mathematical Engineering

Abstract

Proton therapy is an emerging alternative to conventional radiotherapy which is the most common treatment modality for cancer. Proton therapy allows a better sparing of the healthy tissues surrounding the tumor. To plan a treatment, medical physicians and physicists must do a time-consuming work through a specialized software choosing the best trade-off between tumor coverage and sparing of healthy tissues for a given patient.

In recent years, deep learning has been proposed to automate the treatment planning process. However, most existing models only predict one dose distribution which is not always the most optimal for a given patient. An emerging method is Pareto optimization. It consists in allowing the medical physicists to choose the best trade-off for a patient through changing the inputs in a deep learning model.

This master thesis presents a deep learning model for the prediction of a dose distribution with a scaled dose in an organ. The organ to scale and the scaling can be set as an input of the model. The deep learning architecture used is the HD U-Net. The predicted dose distribution is then analyzed and made clinically deliverable through dose mimicking. This thesis focuses on head and neck cancer and the possible late effects due to a proton therapy treatment. The predictions are thus evaluated through the computation of the Normal Tissue Complication Probability (NTCP) of a common complication: xerostomia.

Acknowledgements

First of all, I would like to thank my supervisors, John Lee, Ana Barragan, and Jimmy Olsson. In particular, I would like to thank John Lee for his good advice and innovative ideas, always given with a touch of humor. I am also very grateful for Ana's help and encouragement. I would like to thank her for always being present through the long days of scripting and for being there to answer all my questions.

Then, I would like to sincerely thank Margerie for her continuous help and support throughout the year. I also want to thank her for welcoming me in her office.

Furthermore, I would like to thank Tianfang Zhang for giving us his time to explain his work and Ph.D. thesis in more detail and for sharing some of his code with us. It was very instructive to be immersed in his work and be able to discuss it with him.

Additionally, I am thankful for Elena's help in the scripting in the RayStation, and particularly, I would like to thank her for helping us with the dose mimicking.

Finally, I would like to thank the MIRO lab for their warm welcome. I have enjoyed my time there and I could not have wished for a better place to do my thesis.

Contents

Acronyms	5
1 Introduction	6
2 Context	9
2.1 Radiotherapy Treatment Planning Workflow	9
2.1.1 Imaging the patient	9
2.1.2 Contouring and dose prescription	11
2.1.3 Plan optimization	12
2.1.4 Plan evaluation	13
2.1.5 Treatment planning uncertainties	14
2.2 Proton Therapy	15
2.3 Late side effects of head and neck cancer treatments	16
2.4 Deep Learning	17
3 Literature review on automatic planning	22
3.1 Multi-Criteria Optimization (MCO)	22
3.2 Deep Learning Dose Prediction	24
3.2.1 Conventional dose prediction	24
3.2.2 Pareto dose prediction	27
4 Methods	30
4.1 Patients Database	30
4.2 Artificial Pareto Dose Generation	31
4.3 Network Architecture	33
4.4 Dose prediction models	34
4.4.1 Conventional dose prediction	35
4.4.2 Scaled dose prediction	35
4.4.3 Dose prediction with different scalings in a single organ	35
4.4.4 Pareto dose prediction	36
4.5 Dose mimicking	36
4.6 Normal Tissue Complication Probability	38

5	Results	40
5.1	Dose prediction models	40
5.1.1	Conventional dose prediction	40
5.1.2	Scaled dose prediction	42
5.1.2.1	Model predicting a lower dose in the right parotid	42
5.1.2.2	Model predicting a lower dose in the right submandibular gland	43
5.1.3	Dose prediction with different scalings in a single organ	45
5.1.4	Pareto dose prediction	47
5.2	Dose mimicking	48
5.2.1	Mimicked plan for a lower dose in the right parotid for patient ANON1089	50
5.2.2	Mimicked plan for a lower dose in the right submandibular gland for patient ANON1089	52
5.3	Normal Tissue Complication Probability	54
6	Discussion	56
7	Conclusion	59

Acronyms

ANN	Artificial Neural Network.
CNN	Convolutional Neural Network.
CT	Computed Tomography.
CTV	Clinical Target Volume.
DNN	Deep Neural Network.
DVH	Dose Volume Histogram.
FDG	Fluorodeoxyglucose.
GTV	Gross Target Volume.
IMPT	Intensity-Modulated Proton Therapy.
IMRT	Intensity-Modulated Radiation Therapy.
MCO	Multi-Criteria Optimization.
MRI	Magnetic Resonance Imaging.
MSE	Mean Squared Error.
NTCP	Normal Tissue Complication Probability.
OAR	Organ at Risk.
PET	Positron Emission Tomography.
PTV	Planning Target Volume.
ReLU	Rectified Linear Unit.
ROI	Region of Interest.
SOBP	Spread-out Bragg peak.
VMAT	Volumetric Modulated Arc Therapy.

1 | Introduction

Cancer is one of the major causes of death in the world. In Belgium, cancer was the second leading cause of death in 2019 after heart diseases [1]. The most common cancer treatments are surgery, radiotherapy, and chemotherapy [2]. Radiotherapy is a very effective method but, unfortunately, it also damages the organs surrounding the tumor, also known as Organ at Risk (OAR).

Proton Therapy is an emerging alternative. Unlike conventional radiotherapy, which sends photons to the tumor, proton therapy is a method using protons. Protons are particles characterized by the Bragg peak. As can be seen in Figure 1.1, a proton will deliver most of its energy at the end of its trajectory and nothing after while X-rays have a decreasing energy penetration. The energy delivered by the proton can thus be controlled to target the tumor as precisely as possible and have less damage to the surrounding healthy tissues. Figure 1.2 shows an example of a proton therapy treatment with two beams and a conventional radiotherapy treatment, also called Intensity-Modulated Radiation Therapy (IMRT), with four beams. One can see that the dose after the tumor is nearly zero for the proton therapy case, while a significant dose is delivered to the healthy tissue in the case of the IMRT treatment.

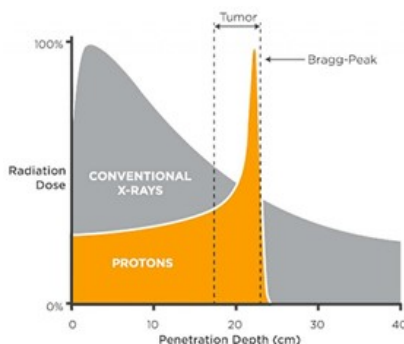


Figure 1.1: Penetration of the radiation dose in proton therapy and conventional radiotherapy. Credits: [3]

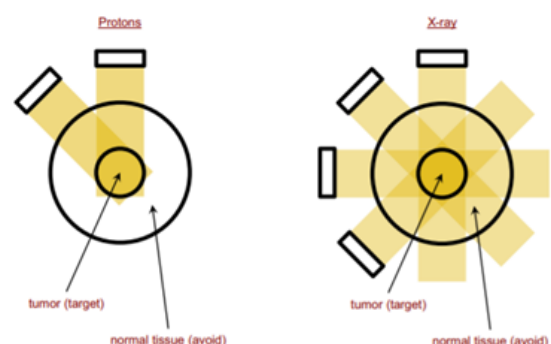


Figure 1.2: Proton therapy has a smaller integral dose than radiotherapy (X-ray IMRT). Credits: [4]

Although proton therapy is a very promising treatment method, it still inevitably affects some OARs. To obtain a treatment plan, medical physicists and physicians have to do a

laborious and time-consuming work, testing different combinations to obtain an optimal treatment plan for the patient. They have to choose between different trade-offs such as tumor coverage and sparing of healthy tissues. This process can take several hours to days of work [5, 6, 7] through a specialized software called treatment planning system. The treatment plan can be subject to inter- and intra-observer variabilities such as the lack of time of the professionals or their skills. This may lead to suboptimal plans [6, 7]. Knowing that the Tumor Volume Doubling Time for a head and neck cancer has a median of 30 days for half of the patients [8], delaying a treatment plan by one week may largely decrease the local tumor control and the patient survival depending on the type of cancer [9].

In recent years, deep learning has started to be used to automate and standardize the process of finding the optimal dose distribution for a given patient’s anatomy (i.e. from a Computed Tomography (CT) and image of organ and tumor volume) [6]. This does not only save time but also avoids the inter-observer variability. However, since the treatment plan must meet different objectives, it is difficult to translate it into a single plan that meets the specifications. A solution can be Pareto-optimization. Pareto-optimality, named after the sociologist Vilfredo Pareto, is a concept that was first presented in economics. When considering the allocation of resources, a Pareto-optimal solution is a state where the situation of one individual cannot be improved without worsening the situation of at least another individual. All Pareto-optimal solutions can then be represented on a continuous Pareto surface [10]. In the same way, the concept can be applied to the choice of radiotherapy treatment accounting for the best trade-off. To achieve the same target coverage, sparing one organ will imply that another organ receives more radiation resulting in different Pareto-optimal solutions. A treatment plan is thus Pareto-optimal if there is no other plan that is better in every way without worsening one of the objectives (more dose in one organ or less dose coverage). As in economy, the set of Pareto-optimal treatment plans can be represented on a Pareto surface. Through the navigation in the set of treatments on the Pareto surface, the physician will be able to choose the best clinical trade-offs with the guarantee of Pareto-optimality. This method is commonly called Multi-Criteria Optimization (MCO). Recently, some deep learning models have been presented where the physician can set a given trade-off as an input of the model. The deep learning model then predicts the corresponding Pareto optimal dose distribution [11, 12, 13].

This thesis aims to adapt an already working deep learning model that predicts an optimal dose distribution to spare the surrounding healthy tissues of the tumor. To do so, the same deep learning model is trained using artificial Pareto dose distributions for the training patients instead of the clinically deliverable doses. These artificial Pareto dose distributions are obtained through a certain scaling of the organ we desire to spare. The most advanced model presented in this thesis is a model allowing to change the dose

by a certain scaling in a given organ where the scaling and the organ can be chosen by setting it as an input of the model. Due to the rather unrealistic artificial Pareto dose distribution used as a ground truth for the models, a mimicking step is done to find the machine parameters allowing us to get a clinically deliverable dose distribution as close as possible to the predicted scaled dose distribution.

We will focus on head and neck cancer. Due to the numerous OARs, radiation in this zone may lead to permanent or late effects such as chronic xerostomia, chronic fibrosis, edema, trismus, dysphagia, or other organ dysfunctions [14]. Xerostomia is one of the most known late effects and we will concentrate on this one only. Xerostomia is a disease affecting the salivary glands and resulting in a dry mouth due to a change in volume, consistency, and pH of secreted saliva [15]. It will therefore be a matter of trying to spare the salivary glands: left parotid gland, right parotid gland, left submandibular gland, and right submandibular gland in the Pareto-optimal treatment plans generated for this thesis. The impact of reducing the dose in one of these organs on the late side effects will be evaluated through the computation of the Normal Tissue Complication Probability (NTCP).

This report will first present the different useful concepts related to radiotherapy, proton therapy, and deep learning in the first chapter. Then, a deeper analysis of the methods used for automatic planning and the generation of different Pareto dose distributions is presented in the literature review chapter. Afterward, the methods used for the experiments of this thesis are presented, followed by the corresponding results. To finish, a discussion of the results and the conclusion are presented.

2 | Context

This chapter aims at explaining the different concepts used in this thesis. First, we will go through the radiotherapy treatment planning workflow. Then, a short introduction to proton therapy will be presented followed by its potential late side effects for head and neck cancer treatments. Finally, an overview of the useful deep learning concepts is presented.

2.1 Radiotherapy Treatment Planning Workflow

Radiotherapy is a cancer treatment method exposing the tumor to ionizing radiation [4]. It can be either in the form of external beam radiation therapy where the patient receives an external x-ray dose or in the form of brachytherapy also called internal beam radiation therapy where a high irradiating seed is placed next to or in the tumor through surgery[16]. The delivered radiation dose is measured in Gray [Gy].

Photon radiotherapy is an external beam radiation therapy using x-rays. It is one of the most common modalities of treatment for cancer today.

The first step in treating a cancer patient is treatment planning. The planning of a treatment with radiotherapy follows different steps. First, it is needed to have an accurate image of the patient. Once this is done, the organs at risk and the tumor are contoured. Then follows the optimization of the dose to meet certain objectives such as reaching the minimum dose on the tumor and reducing the dose in the organs at risk. Finally, the treatment plan is evaluated to validate it or not. If it does not meet the specifications required, a new optimization is done. Throughout these steps, dosimetrists, medical physicians and physicists work together to deliver the most valuable plan.

2.1.1 Imaging the patient

Imaging the patient is the very first step in planning a treatment. The Computed Tomography (CT) scan is obtained with an X-ray scanner. A gantry is rotating around the patient lying on a table, as shown in Figure 2.1. The patient might wear an immobilizing mask, as shown in Figure 2.2, to prevent any motion resulting in an incorrect CT scan.

Thanks to the gantry, the machine produces cross-sectional CT scan slices. Each slice will represent a thickness of the tissue between 1 and 10 millimeters [17]. Once the desired number of slices has been generated, a 3D image can be reconstructed. This 3D image gives the radiation attenuation coefficient due to different densities in the body. This radiation attenuation coefficient is given in Hounsfield Units (HU). This will be converted in proton stopping power ratio (SPR) to be able to compute the proton therapy dose [5].



Figure 2.1: CT scanner disposition.
Credits: Terese Winslow [17]

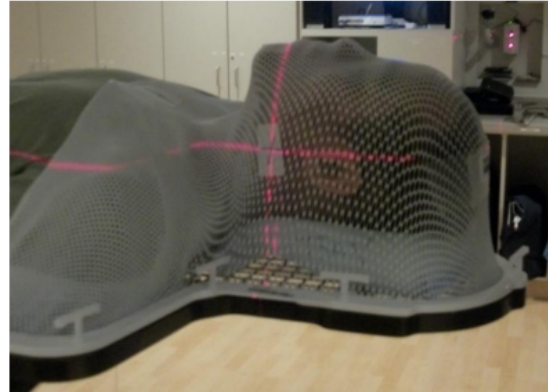


Figure 2.2: Thermoplastic head and neck mask to immobilize the patient.
Credits: [18]

Usually, the computed tomography is not sufficient for delineating the OARs and the tumor volume, also called Gross Target Volume (GTV). A CT scan is often combined to a Magnetic Resonance Imaging (MRI) or Positron Emission Tomography (PET) scan to improve visualization and contouring of a Region of Interest (ROI) [5, 18]. Although MRI can be helpful, PET scan is the most used imaging method in addition to CT scan [5]. A PET scan is obtained by doing a CT scan while injecting a radioactive tracer to the patient [18]. The most used tracer is Fluorodeoxyglucose (FDG) but other tracers might be used to counteract the limitations of FDG such as FMISO or Cu-ATSM [5, 18].

The axial, sagittal and coronal imaging with the four different methods of imaging can be observed in Figure 2.3. The use of a CT scan together with a PET scan is considerably increasing the visualization. The contours of the ROIs can then be drawn.

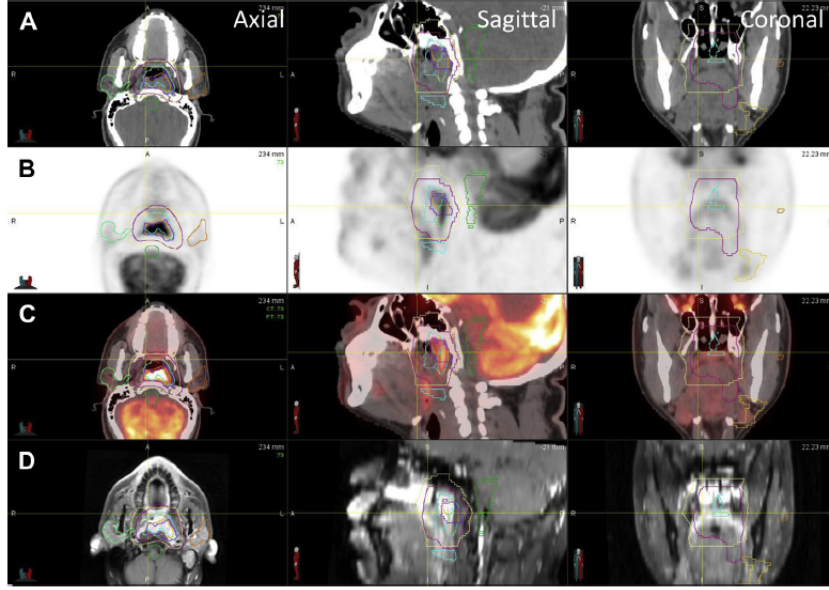


Figure 2.3: (A)CT scan (B)PET scan (C)PET/CT (D)MRI. Credits:[18]

2.1.2 Contouring and dose prescription

The contouring of the ROIs is usually done by a radiation oncologist thanks to a CT scan and often using auxiliary information from other imaging modalities like MRI or PET scans. Some research has shown that contouring using CT in combination with another imaging modality, instead of CT alone, shows better results [18]. The target volume to be irradiated is contoured following a three-step process: the Gross Target Volume (GTV), the Clinical Target Volume (CTV) and the Planning Target Volume (PTV). The tumor visible through imaging is the GTV. To account for microscopic and suspected tumor sites, the GTV is enlarged by 5 to 15 mm to result in the CTV [18]. To account for uncertainties in the radiotherapy treatment planning, an additional expansion of 3 to 5 mm to the CTV can be made to create the PTV [18]. An illustration is shown in Figure 2.4. Note that for proton therapy treatments, instead of using a PTV, the plan is created directly using the CTV volume, using advanced techniques for robust optimization to account for uncertainties (See Section 2.2).

In addition, the contouring of the OARs can be different depending on the radiation oncologist. The inter-observer variability can thus be very high [19]. To remediate this, some consensus guidelines have been set up [19, 20]. The consensus for contouring the OARs for head and neck patients can be seen in Figure 2.5. The different numbers on the figure account for the following organs: 1) parotid glands, 2) pharyngeal constrictor muscles, 3) carotid arteries, 4) spinal cord, 5) mandible, 6) extended oral cavity, 7) buccal mucosa, 8) lips, 9) brain, 10) chiasm, 11) pituitary gland, 12) brainstem, 13) supraglottic larynx, 14) glottic area, 15) cricopharyngeal inlet, 16) cervical esophagus and 17) thyroid [19].

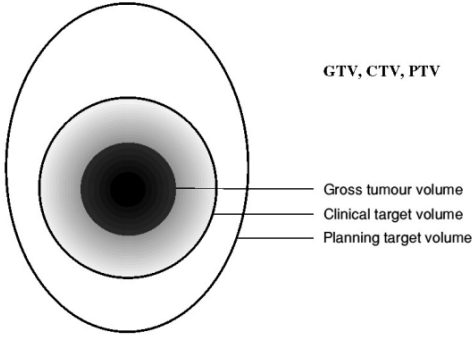


Figure 2.4: Contouring of the target volume. Credits: [21]

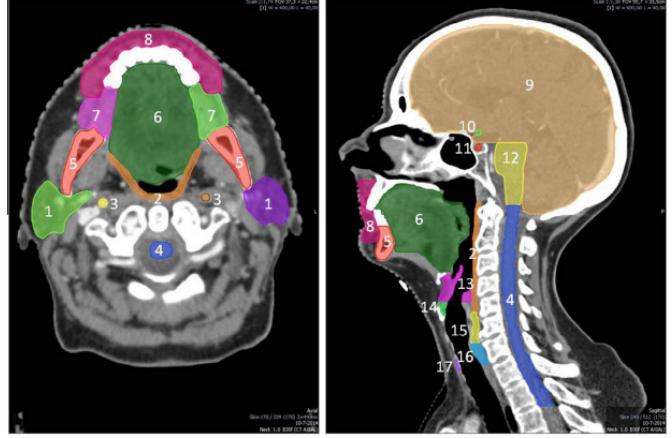


Figure 2.5: Contouring of head and neck OARs. Credits: [19]

Once the contours of the ROIs are drawn, the dose prescription for each ROI can be set. A minimal dose prescription is set for the tumor and a maximal dose prescription is set for the OARs depending on their dose tolerance [18].

2.1.3 Plan optimization

From the CT scan, the contours and the dose prescriptions, a treatment plan can be generated. The maximal and minimal dose prescriptions give different constraints and objectives and the problem can be rewritten in a Multi-Criteria Optimization (MCO) problem. A possible mathematical formulation of the optimization problem is given by Equation 2.1 [22, 23].

$$\begin{aligned}
 & \underset{x}{\text{minimize}} && \{f_{PTV}(x), f_{OAR_1}(x), f_{OAR_2}(x), \dots, f_{OAR_n}(x)\} \\
 & \text{subject to} && x \geq 0 \\
 & && c_j(x) \leq 0, \quad j = 1, \dots, m
 \end{aligned} \tag{2.1}$$

The patient geometry is discretized using voxels. The $f_{ROI}(x)$ are the individualized objective functions for each ROI, $c_j(x)$ are the different constraints related to the ROIs and the patient geometry and x is a vector of weights (or intensities) for the elementary dose beamlets [24]. The vector x is typically referred to as "fluence map intensities" in conventional radiotherapy [11] and "spot weights" in proton therapy [22].

$$\begin{aligned}
 & \underset{x}{\text{minimize}} && \sum_{i=1}^n w_i f_i(x) \\
 & \text{subject to} && x \geq 0 \\
 & && c_j(x) \leq 0, \quad j = 1, \dots, m
 \end{aligned} \tag{2.2}$$

A scalarization of the problem is necessary to be able to solve it. The most common method is a weighted sum of the different objective functions associated with the differ-

ent ROIs. This optimization problem is given by Equation 2.2. The weights must be non-negative and are given by w_i for each ROI i .

Changing the weight of a contour in the objective function from Equation 2.2 will result in a different optimal treatment plan. This allows the dosimetrist (in cooperation with a physician) to choose between different trade-offs iteratively. In the MIRO laboratory, the RayStation software (RaySearch, Stockholm, Sweden) is used to solve this type of optimization problem. This software generates a dose map satisfying the objectives by changing the machine parameters such as the proton beam intensities. This process of changing the machine parameters to obtain a certain dose distribution is called inverse planning.

A treatment plan is called a Pareto optimal plan if decreasing the value of one of the objective functions in Equation 2.1 implies increasing another objective function. Thus x^* is a Pareto optimal solution if there exist no other plan that decreases $f_i(x^*)$ without increasing the other objectives $f_j(x^*)$ with $j \neq i$. Each Pareto plan corresponds to a certain combination of weights in Equation 2.2. It is commonly said that the set of all Pareto plans spans a Pareto surface.

In recent years, automation of the treatment plan optimization step is a major research subject (see Chapter 3). Today, most automated methods use deep learning to predict only one optimal dose map based on previous patients. Then, dose mimicking is used to try to reproduce the same dose map by adjusting the machine parameters. This considerably decreases the treatment planning duration. But, in radiotherapy, there is an infinity of optimal doses located on a Pareto surface. Limiting the dose in a certain region means that there will be an increase in the dose in another region. In conventional treatment planning, the dosimetrist (and a physician) must choose the best trade-off for healthy tissue sparing. This thesis tries to address this problem by proposing a deep learning model with the possibility of providing which OAR to spare as an input.

2.1.4 Plan evaluation

During the treatment planning optimization process, the quality of the generated plan must be evaluated. This is done by a Dose Volume Histogram (DVH) for each contour. It depicts the percentage of the volume that receives at least a given dose in Gray [Gy]. An example of a DVH graph is represented in Figure 2.6. It is useful when needing to have a general overview. If there is a need to compare values precisely, it is possible to use the dose volume metrics which are in fact points on the graph. Two types of dose volume metrics are commonly used:

- D_x : The highest dose that at least x % of the volume receives.

- V_x : The greatest percentage of the volume that receives a dose of at least x Gy.

By convention, the maximal dose a ROI receives is given by D_2 or D_5 and the minimal dose is given by D_{98} or D_{95} . This avoids considering a single voxel that has a lower or higher value than the rest of the ROI volume. The D_{mean} is used when the average dose in a ROI has to be quantified.

In this thesis, we will consider D_{95} to compare the minimal dose levels in the target volumes and the D_{mean} to evaluate the average dose in an OAR.

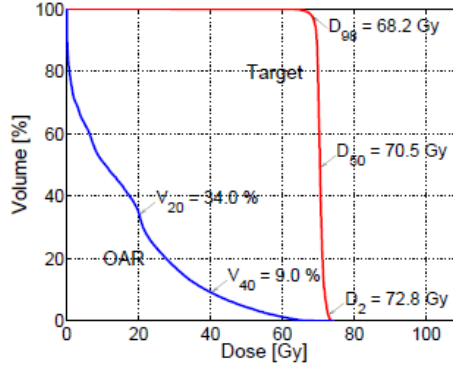


Figure 2.6: Example of DVH curves for the target (red) and an OAR (blue). Credits: [16]

2.1.5 Treatment planning uncertainties

In the treatment planning workflow, the dosimetrist also needs to account for treatment planning uncertainties and aim at robustness. Two types of errors can be identified: systematic and random errors. The treatment is typically delivered over a couple of days to a couple of weeks. The treatment delivery is divided into fractions with usually one fraction per day. Systematic errors are the same for every fraction of treatment, while random errors vary from fraction to fraction. Both types of errors cause a change in dose distribution, sending less dose to the tumor and more to the surrounding healthy tissue, or the opposite behavior [5].

Typical sources of uncertainties are geometric such as an inaccurate CT scan, the motion of the patient during treatment planning and the motion of the organs even if the patient is immobilized [25]. To account for these uncertainties in radiotherapy, a margin on the CTV is taken. The PTV is used instead to contour the target volume. The static PTV then accounts for the motion of the CTV [18]. A similar margin can be applied to OARs to result in the planning OAR volume also denoted as PRV [25]. In proton therapy, the CTV is used instead of the PTV and robust optimization is used to account for errors (See Section 2.2).

2.2 Proton Therapy

In recent years, proton therapy is emerging as an interesting alternative to photon radiotherapy. Proton therapy is an external beam radiation therapy sending protons instead of photons to the tumor. Protons are charged particles that can be taken from hydrogen gas. They are heavy and can be accelerated to a certain fraction of the speed of light using an accelerator such as a cyclotron or a synchrotron [4]. Proton therapy presents a lower integral dose than radiotherapy as shown in Figure 1.2. Despite its cost, it is a very promising method for cancers located close to critical healthy tissues such as the brain, the lungs or the heart. It is also used for children vulnerable to late side effects [4] and in case of re-irradiation [4, 26].

Its lower integral dose is due to a property of the proton called Bragg peak. The proton will deliver most of its energy at the Bragg peak with almost no exit dose. This results in a steep dose gradient at the Bragg peak [25] while X-rays have a decreasing emission of energy as depicted in Figure 1.1. The location of the Bragg peak of a proton is determined by its weight and the density of the tissues it goes through. But a simple Bragg peak, as shown in Figure 2.7a, is not enough to cover the tumor. By superposing Bragg peaks of protons of increasing weights, a beam presenting a flat Spread-out Bragg peak (SOBP), as shown in Figure 2.7b, can be obtained. Two different techniques are used to broaden the narrow proton beam. The first method is passive scattering, which uses different pieces of material to scatter the beam and conform it to the shape of the tumor [4, 25]. Another method, which is now the state of the art, is pencil beam scanning. This technique uses magnets to deflect the very fine beams to scan the tumor slice by slice [25].

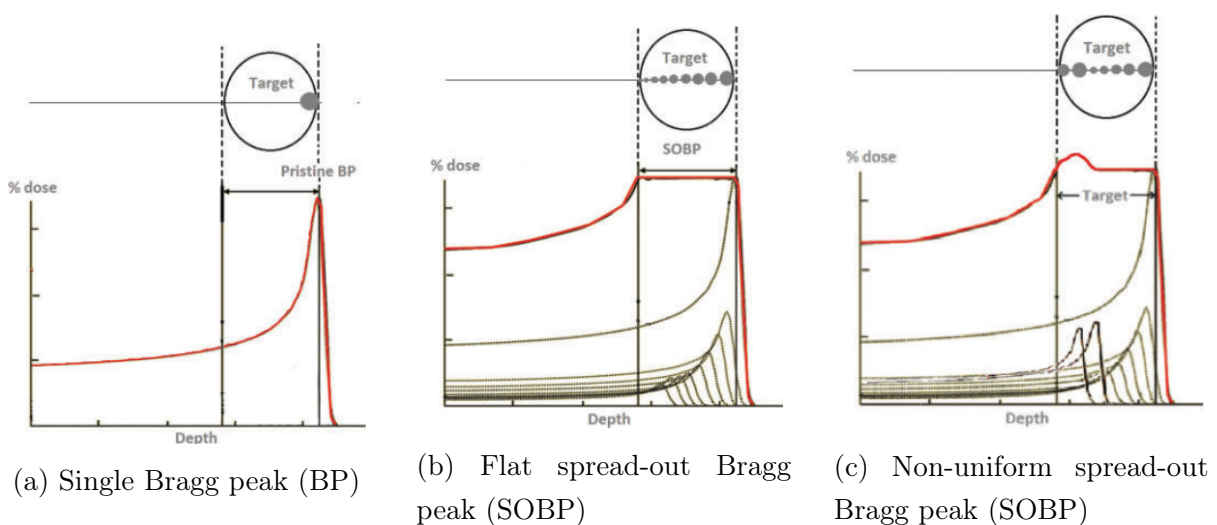


Figure 2.7: Schematic overview of Bragg peaks. Credits: [25]

Additionally, the way the dose is delivered by each beam can be different. First, there

is the single field uniform dose (SFUD) where each beam delivers a uniform dose [25, 26]. Moreover, there is multi-field optimization that optimizes all the beams at the same time. Each beam can deliver a different non-uniform dose, but the final dose is homogeneous [26]. The fact that we are able to modulate the intensity of each beam is called Intensity-Modulated Proton Therapy (IMPT)[25, 26]. IMPT can only be achieved when using pencil beam scanning [25]. When the intensity of a beam is modulated to be non-uniform, a non-uniform SOBP is used [25]. The non-uniform SOBP is obtained by superposing protons of different weights [25] as shown in Figure 2.7c. IMPT is the method used in this thesis.

The treatment planning workflow for proton therapy is the same as for conventional radiotherapy. However, the planner must account for additional uncertainties due to the properties of the proton. The position of the Bragg peak is difficult to predict and depends on the density of the tissues the proton goes through. All the factors that may lead to an error in the calculation of the position of the Bragg peak are called "proton range uncertainties". An inaccurate CT scan resulting in inaccurate tissue densities might provoke a misplaced Bragg peak. This may lead to an exposition of too much radiation for the healthy tissues and not enough radiation on the tumor [26]. An inaccurate CT scan can be due to the patient's motion, a change in patient weight throughout the treatment or an increase in the tumor volume [25, 26]. Additionally, a bad calibrated CT can result in an inaccurate measurement of the densities of tissues.

In the case of proton therapy, using a margin on the CTV is not working anymore to account for uncertainties [25]. Reimaging the patient regularly during the treatment may already reduce uncertainties [26]. To account for these uncertainties in the treatment planning, robust optimization is used. Usually, 21 scenarios are elaborated for the different shifts in the three space dimensions and different proton range errors (3 scenarios accounting for proton range errors and 7 to account for systematic setup error). The scenario having the worst objective function is optimized. This is called a worst-case scenario optimization or minimax [27].

2.3 Late side effects of head and neck cancer treatments

Despite its lower integral dose, proton therapy still leads to toxicities in the tissues in the target range of the proton beam [26]. Due to the high number of OARs for head and neck cancers and the closeness of the tumor and the OARs, there is still a high risk of side effects due to proton therapy treatments [28].

It is typically assessed by the Normal Tissue Complication Probability (NTCP). Many models aim at predicting the NTCP for head and neck cancer patients depending on the

toxicity they receive [28]. These models help the dosimetrists and the physician choose the best trade-off between target coverage and organ sparing. Xerostomia and sticky saliva are reported as the most common complications and they result from radiation in the parotid glands, submandibular glands and sublingual glands [28]. Xerostomia results in a dry mouth due to a change in volume, consistency, and pH of secreted saliva [15].

2.4 Deep Learning

Machine learning is a branch of artificial intelligence where a computer algorithm mimics the learning of a human [29]. It typically works in two steps: training and inference [29]. Based on the analysis of training data, the algorithm learns to detect patterns automatically [30]. After the training step comes the inference step. Based on new data, the algorithm will make a prediction, make a decision or perform classification with minimal intervention of humans [29, 30].

Machine Learning can be divided into two complementary categories: supervised and unsupervised learning. **Supervised learning** is when the training data contains labeled [input, output] pairs. The model is then trained to generate the corresponding output to a certain input [29]. When the desired output is a continuous value (for example, a dose map), it is called regression and, when it comes to predicting a categorized value, we talk about classification [30]. On the contrary, **unsupervised learning** is working with unlabelled inputs and does not consider output data. It aims at discovering common patterns in the data [29, 30]. Some applications are outlier detection and clustering [29].

Today, supervised learning is the most used method in the medical imaging sector. Research is underway to use semi-supervised learning or self-supervised learning which is between supervised and unsupervised learning. The outputs would then partially be known. It could reduce errors and take less time to label the data [29].

As described above, regression is a common supervised task. Linear regression is the most known mathematical model but other more specific models exist to approximate an output from multiple inputs. The most commonly used are Artificial Neural Networks (ANNs) [29]. This model is inspired by the connection of neurons in the human brain [30]. A neuron consists of dendrites, a nucleus and an axon. Multiple inputs enter the neuron through the dendrites, the cell or nucleus processes the information and the output goes through the axon to the synapse that transmits the information to the dendrites of the next neuron. The artificial neuron is based on the same structure. Multiple inputs are processed in the nucleus with a nonlinear activation function to result in a single output as can be observed in Figure 2.8 - left. The neurons are then combined in a network with possibly hidden layers (see Figure 2.8 - right). An ANN is thus a network of intercon-

nected artificial neurons. When using multiple hidden layers, it is called a Deep Neural Network (DNN) [31].

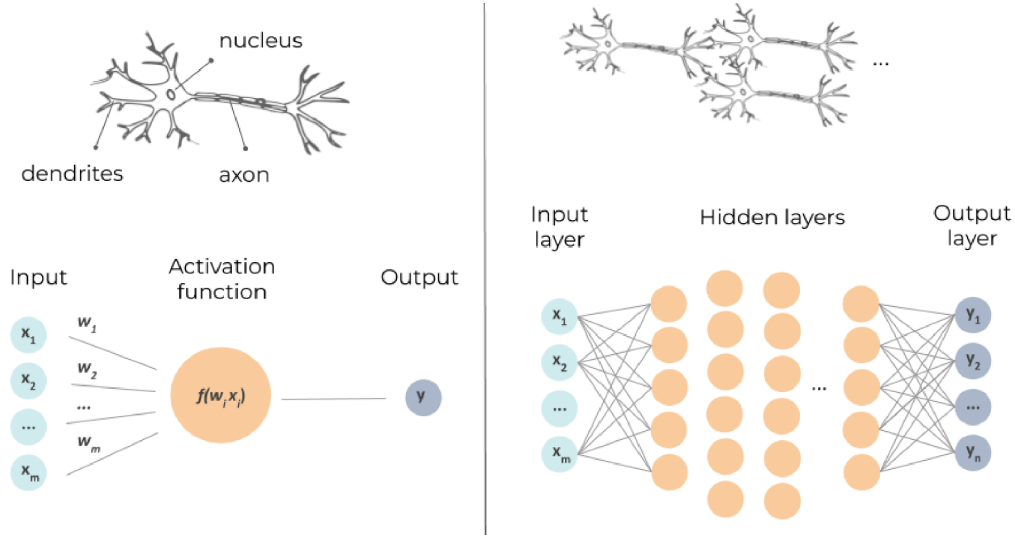


Figure 2.8: Comparison between human neuron and artificial neuron (left) and comparison between human and artificial (deep) neural network (right). Credits: [29]

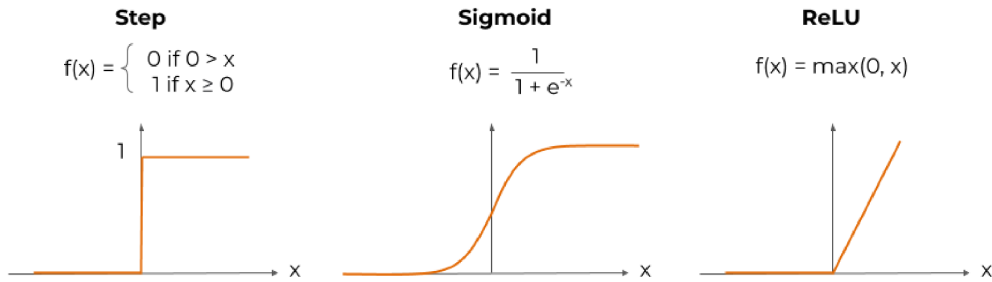


Figure 2.9: Common activation functions. Credits: [29]

Each node of an ANN will have an output based on an activation function applied to weighted inputs [32]. Common activation functions are represented in Figure 2.9. The aim of an ANN is to minimize the error between the prediction in the output layer and the ground truth. This is done by adjusting the weights of each node based on the error computed at each iteration of the training [31]. The error is often defined by a loss function that has to be minimized through the training process. The most common loss function is the Mean Squared Error (MSE).

For the training of the model, the dataset is divided into a training (generally large) and validation (generally small) dataset. One iteration is called an epoch. In one epoch, the model goes over the whole training dataset. In a large dataset, the model has to be updated more often than after seeing all the patients. That is where a batch size has to be determined. After each batch, the model weights will be updated and the output is

evaluated. As long as the model is not performing well on the training patients, we have underfitting and the training can be continued. At the end of each epoch, the output of the model is evaluated with the validation dataset. If the validation loss is high while the training loss is low, the model is too specific to the training patients and we have overfitting [31].

Deep learning is a subclass of machine learning where a DNN is used. An adaptation of a DNN to deal well with images is a Convolutional Neural Network (CNN). A CNN mimics the visual working of an animal cortex [32]. It is typically composed of convolution, pooling, and fully connected layers [33] as shown in Figure 2.10.

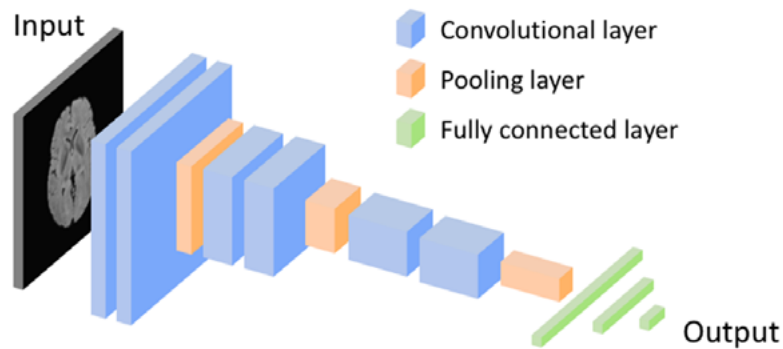


Figure 2.10: Schematic representation of a Convolutional Neural Network (CNN). Credits: [32]

A **convolutional layer** allows to extract features from an image [30, 32, 33] and to detect local patterns and motifs in an input image [30]. The discrete convolution is a linear operation that applies a small array of numbers, called a kernel, across the input array. It is an element-wise product of the kernel pixels and a block of the input array (of the same size as the kernel), which is then summed up to result in a single numerical value in the feature map as depicted in Figure 2.11. The kernel size in 2D is typically of size 3x3 but might be 5x5 or 7x7 [31, 33]. At each step of the convolution, the center of the same kernel is translated over the input array by a certain number of pixels, this is called the stride. Typically, a stride of 1 is used but a greater stride might be chosen to have a dimension reduction in the feature map.

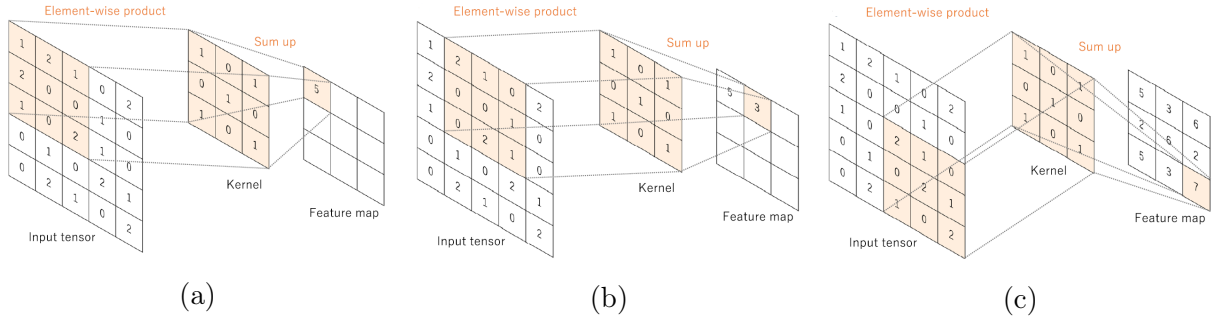


Figure 2.11: Operation of convolution on a small example for a stride of 1, a kernel of size 3x3 and no padding. Credits: [33]

But, the convolution operation, as described above, does not allow placing the center of the kernel at the extreme pixels of the input array. This results in decreasing the dimensions of the feature map. To remedy this, padding is used. This is most often done by adding rows and columns of zeros at the sides of the array. The padding technique is represented in Figure 2.12. Once the feature map is obtained, it goes through the activation function, typically the Rectified Linear Unit (ReLU) which is the most popular one [7] presented in Figure 2.9. Multiple convolution layers can be used in a CNN. In each of them, multiple convolution neurons can be used. It implies that, for the same input array, multiple kernels are used to obtain different feature maps [31]. Each of them will then show a different characteristic of the input array [33]. Training a convolutional layer implies finding the optimal kernels based on the training dataset. The number of convolutional layers, the number of convolutional neurons in each layer, the stride, the padding and the size of the kernel are hyperparameters chosen beforehand [33].

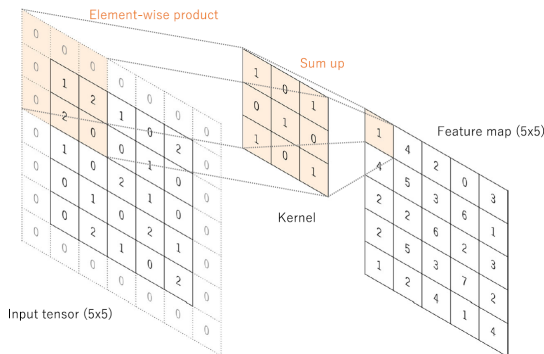


Figure 2.12: Example of a convolution operation using zero padding with a kernel size 3x3 and a stride of 1. Credits: [33]

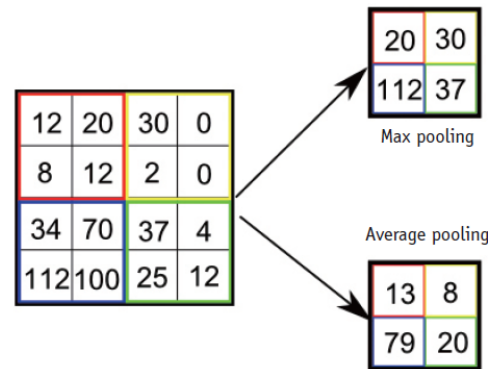


Figure 2.13: Example of max and average pooling operations. Credits: [30]

A **pooling layer** allows reducing the dimension of the feature map. This will reduce the requirement in computational resources to have the right output [31] but this will also make the convolutional layers less sensitive to small shifts and distortions in the image

[30, 33, 34]. The most common types of pooling are max pooling and average pooling. It is typically using a 2x2 filter with a stride of 2 to reduce the feature map dimensions by a factor of 2 [31, 33]. As represented in Figure 2.13, the max pooling will take the maximum element covered by the filter and the average pooling will take the average of all elements covered by the filter. Usually, the training does not fix any parameter of the pooling layer. The hyperparameters are the size of the filter, the stride, the type of pooling and the padding [33].

At the end of the sequence of convolutional and pooling layers, the features are flattened in a one-dimensional array [31, 33]. Then it enters the **fully connected layers**, also known as dense layers [33], which is a generic neural network [31]. The fully connected layers will produce the desired classification output corresponding to the image. As for a generic neural network, it associates the inputs with a learnable weight that sum up to the desired output [33]. Each fully connected layer is followed by an activation function such as ReLU. During the training, the weights of the fully connected layers are determined. The hyperparameters are the number of weights and the activation function [33].

A CNN performs mainly classification from images, where the output is a single class label [35]. For images to images tasks, fully convolutional networks have been proposed [29]. A fully convolutional network is a CNN without the fully connected layers. It consists of convolutional layers and downsampling and/or upsampling operations. The most used fully convolutional network is U-Net [29] and will be presented in the next chapter.

3 | Literature review on automatic planning

Automatic planning can be achieved in many different ways. First, we will review the different already existing methods to generate a Pareto optimal plan without deep learning, commonly known as MCO. Then, we will discuss the different existing deep learning models for a unique dose prediction and a Pareto dose prediction.

3.1 Multi-Criteria Optimization (MCO)

MCO is aiming at solving the optimization problem depicted in Equation 2.1. The automatic planning framework for MCO without machine learning often proceeds in two steps [16, 22, 24]:

1. The generation of a discrete approximation of the Pareto surface.
2. The automated navigation in the Pareto set to seek the best trade-off for a given patient.

First, the Equation 2.1 must be scalarized (changing the vector of objectives to a scalar formulation using weights) to be able to solve it. Once it is scalarized, the Pareto plan corresponding to each scalarization parameter value can be computed. The most known scalarization method is a weighted sum as depicted in Equation 2.2. Other linear convex formulations exist and are commonly solved through a linear programming solving method such as the Chambolle-Pock algorithm [36]. To accelerate the resolution of linear convex problems, a projection algorithm can be used. A projection algorithm iteratively projects the current solution on the violated convex constraints and is therefore guaranteed to converge to the optimal solution in a finite number of iterations [24]. Such linear scalarizations are shown to work well. But, they do not allow finding accurately plans linked to a non-convex optimization problem while the objective functions are often DVH-based and thus have a non-convex behavior [22]. Alternative methods of scalarization have been proposed. In [22], it has been proposed to use the augmented weighted Chebychev metric (AWCM) to scalarize the equation.

Once the formulation of the problem is scalarized, the Pareto surface can be approximated. Each Pareto plan corresponds to a certain combination of scalarization parameter values. However, choosing uniformly distributed weights does not guarantee that the Pareto points on the Pareto surface are well distributed. Different algorithms have been proposed to choose the optimal combination of parameter values [16]. For example, a Differential Evolution algorithm can be used in combination with a certain metric to evaluate the quality of the Pareto space approximation [22]. Then, the pre-computed points on the Pareto surface can be interpolated to result in a continuous Pareto surface [16].

Finally, when the Pareto surface is obtained, an interface is necessary to be able to navigate through all the plans to find the best trade-off. Typically, in a treatment planning system, the navigation is done using sliders. Moving the slider related to the dose in one OAR will move the other sliders and by consequence change the chosen MCO optimal treatment plan [37]. This allows the medical physician and the dosimetrist to reduce the duration of the time-consuming replanning process to find the best trade-off for a given patient [16]. An example of MCO interface in the software RayStation (RaySearch Laboratories, Stockholm, Sweden) is shown in Figure 3.1.

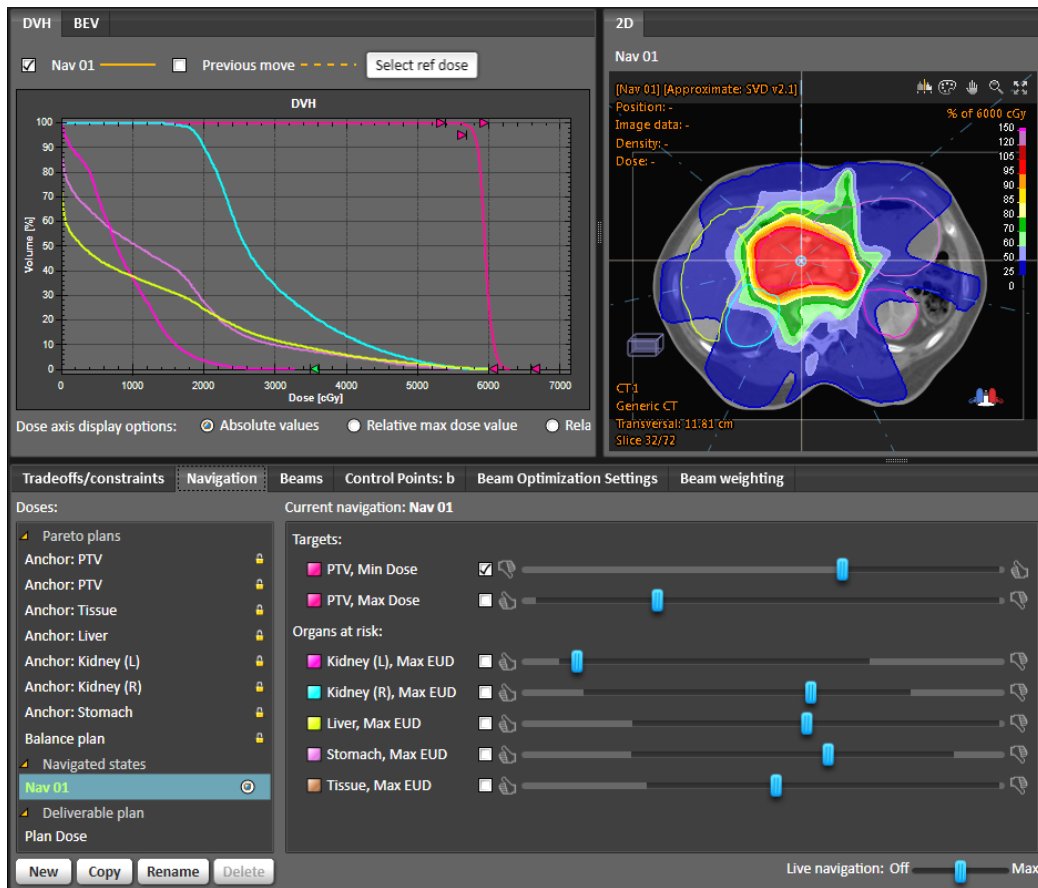


Figure 3.1: MCO Pareto set navigation interface in the software RayStation (RaySearch Laboratories, Stockholm, Sweden) with a representation of the different sliders for PTV and OARs and their corresponding DVH curves. Credits: [16]

3.2 Deep Learning Dose Prediction

3.2.1 Conventional dose prediction

In recent years, multiple deep learning networks have been studied to automate the human part of the treatment planning process (e.g. the choice of weights of the scalarization in the treatment planning system). The most common way of doing so is to predict the dose distribution of the treatment with the patient's anatomy as input to the model. The patient's anatomy commonly includes the CT, the masks of the OARs and the mask of the PTV or CTV, depending on whether it is proton therapy or radiotherapy [6].

To predict the dose distribution with deep learning from input images, an image-to-image model such as a fully convolutional network is necessary [29]. The most popular fully convolutional network is U-Net. It starts with a contracting path presenting a sequence of convolutional and pooling layers. Then, the expansive path presents upsampling and convolutional operations to increase the resolution of the output again [35]. The U-Net upsample operation is the combination of upsampling, convolution and ReLU operations followed by a concatenation of the features set on the other side of the "U" (through skip-connections) [7]. Sometimes, the upsampling and convolution operations are replaced with a transposed convolution. A typical 2-dimensional U-Net architecture can be observed in Figure 3.2.

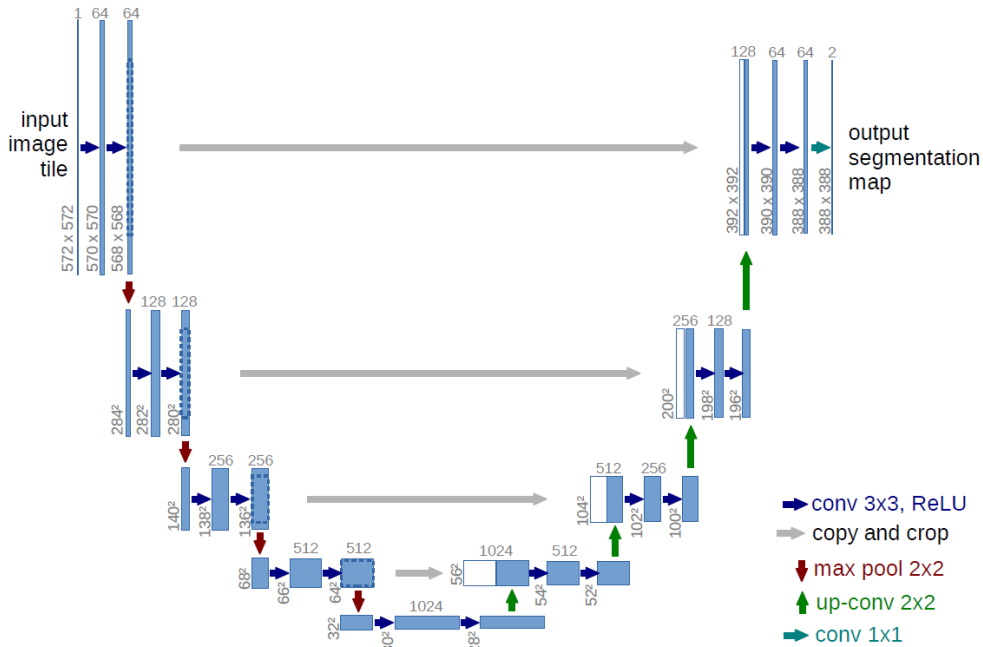


Figure 3.2: Two-dimensional U-Net architecture. The blue boxes represent multi-channel feature maps. The number of channels is written on top of the boxes and their dimensions is written at the bottom left side. The white boxes are copied features and the arrows represent the different operations. Credits: [35]

Here the contracting path presents 3x3 convolutional layers and 2x2 pooling layers. Thanks to the skip-connections between the contracting and the expansive path, the U-Net can take both local and global features into account when doing a prediction [6, 7, 11]. The contracting and expansive paths being symmetric and linked through skip-connections, they form a U-shape with the bottleneck part at the bottom [35].

It is possible to predict the dose in a slice-by-slice manner using a 2-dimensional U-Net [38] as the one presented in Figure 3.2. But, it may result in some errors in the superior and inferior slices of the PTV which might be avoided using a 3-dimensional model [6, 7]. The Standard 3D U-Net is represented in Figure 3.3.

Another commonly used model is the Densely Connected CNN also called DenseNet. DenseNet does not use any pooling layers but densely connects its convolutional layers inside a dense block. The learned feature maps of all previous layers are copied to be used in the current layer as an input of the convolution. This operation is called dense convolution [6, 7]. The structure of the DenseNet can be observed in Figure 3.3. The dense convolution presents some interesting properties: it promotes feature propagation, it decreases the number of trainable parameters and it reduces the vanishing gradient issue. However, even if it is more efficient, it needs too much memory compared to U-Net [6, 7].

A third model is the Hierarchically Dense U-Net, also called HD U-Net. It combines the Standard U-Net with the DenseNet. It takes the same structure as the U-Net but uses dense convolutions and dense downsamplings and keeps the U-Net upsampling operations [6, 7]. The dense convolution consists in convolution and ReLU followed by a concatenation with the previously computed feature maps. The dense downsampling is a strided convolution together with a ReLU that computes a feature set that has half of the resolution of the previous feature set. The computed feature set is then concatenated with the max-pooled previous feature set. The architecture of the HD U-Net can be observed in Figure 3.3. By its U-shape and no dense upsampling operations, the HD-Unet can use less memory than the DenseNet while keeping its interesting properties [7].

A comparison of the performances of these three models for the prediction of the radiotherapy dose prediction for head and neck cancers has been done by [7]. From their research, it seems that HD U-Net outperforms the two other models on the test dataset by taking the advantages of the two of them. First, HD U-Net has 12 times fewer trainable parameters than the U-Net and its training is 4 times faster than the DenseNet. In terms of accuracy, HD U-Net also outperforms the two other networks. Even if it has around the same training loss as the U-Net, it seems to have a lower validation loss and thus generalizes better and there is less overfitting. Despite its interesting properties, DenseNet

performs the worst on all levels [7].

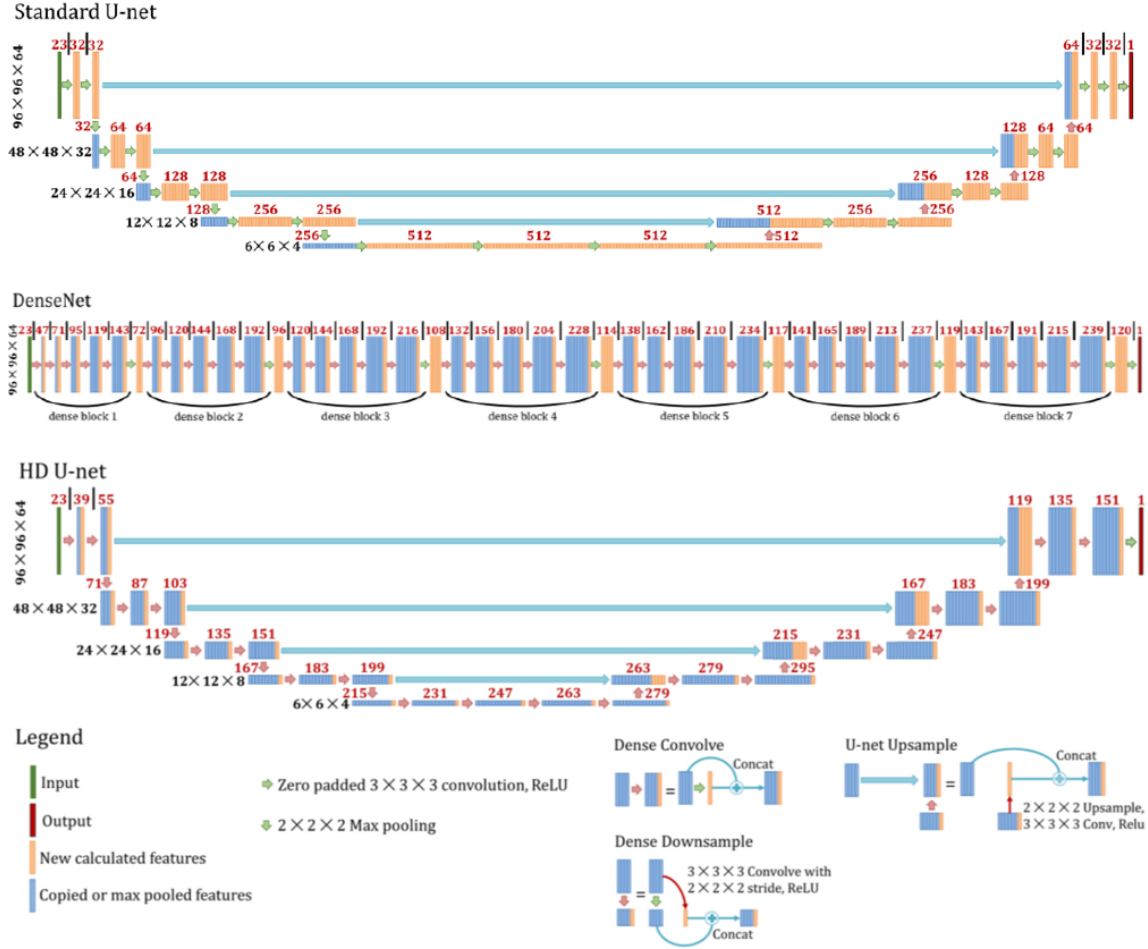


Figure 3.3: Three-dimensional U-Net, DenseNet and HD U-net. The red numbers above the boxes represent the number of feature maps and the black numbers on the left represent the dimensions of each feature map. Credits: [7]

Since the beam configuration can be different depending on the patient, the planner or the institution, it is also interesting to have a model taking into account different beam configurations. To do so, [6] has proposed to include one input channel for beam configuration through a 3D matrix. It has been shown that, for various beam configurations, such a model outperforms the basic model taking only the patient's anatomy in input [6].

Once the dose distribution has been predicted, it has to be transformed into a clinically deliverable plan with the corresponding machine parameters. This is often done through inverse planning by dose mimicking (voxel-wise or through isodoses) or by taking the DVH-based dose objectives from the predicted dose distribution. As an alternative, a second CNN can be used to predict the fluence map for each beam as presented in [38].

3.2.2 Pareto dose prediction

In the previous section, we have presented the prediction of one optimal dose distribution for a specific patient using deep learning. Typically, these models average the trade-offs between target coverage and OAR sparing, present in the training dataset [13]. However, there exist a multitude of optimal plans satisfying different trade-offs. To allow the physician to choose the best trade-off for a given patient, some Pareto frameworks using deep learning have been developed.

To our knowledge, all existing models for Pareto dose prediction are focusing on the prostate cancer treatment with radiotherapy (IMRT [11, 13, 39, 40] or Volumetric Modulated Arc Therapy (VMAT)) [12, 41]). The trade-offs for this type of cancer treatment are between sparing the OARs (i.e., the rectum, the bladder, the left and right femoral heads) and covering the target [11].

The model presented by Dan Nguyen in [11] aims at predicting the Pareto dose distribution linked to a certain trade-off as input of the model. The given trade-off choice is represented by the weights of the scalarization of the MCO problem by a weighted sum (see Equation 2.2). These weights for each ROI are set as input of the model along with the patient's anatomy. For a different combination of weights, a different Pareto optimal dose will be predicted by the model. The Pareto dose prediction workflow is presented in Figure 3.4.

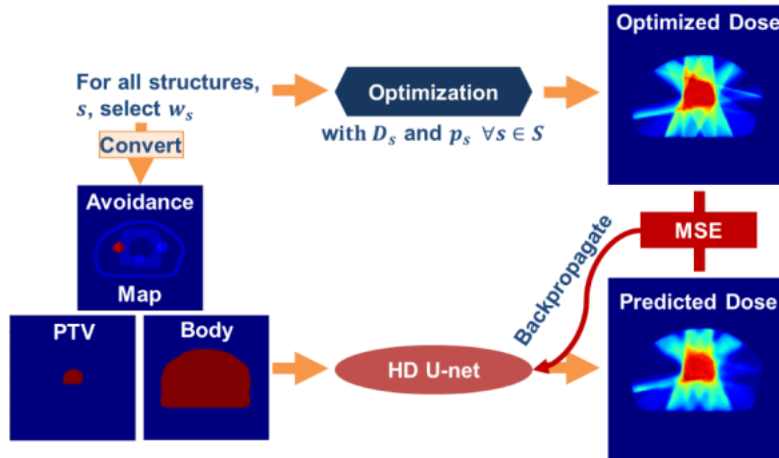


Figure 3.4: Pareto dose prediction workflow by Dan Nguyen et al. [11]. Credits: [11]

The model architecture is the HD U-Net presented in Figure 3.3. As in the previous section, the patient's anatomy is set as input of the model and the weights are included in the different masks of the OARs. The OARs masks are 3D matrices containing 1 if the voxel is in the corresponding OAR and 0 otherwise. In this Pareto model, the ones in the masks are replaced with the corresponding weights of the OAR objective. To be able to train such a model, a great number of plans are necessary to span the Pareto

surface of the training patients. In [11], 1200 plans per patient are generated through the resolution of the weighted sum equation (Equation 2.2) using Chambolle-Pock [36] for different pseudo-random weights. Only clinically acceptable plans were included in these 1200 plans per patient [11].

Due to the good results of the model presented in [11], some additional research has been done on this model with some variations.

First, the influence of additional losses on the MSE has been investigated [39]. Although most deep learning models only use a voxel-wise MSE, the addition of a DVH loss and an adversarial loss has been investigated. Since the DVH is the most used metric to evaluate a treatment plan, it is relevant to include a loss associated with the DVH. Next to the DVH loss, the adversarial loss acts as a discriminator learning its loss. The adversarial loss is there to help the model to further differentiate the small nuances of the different trade-offs that are difficult to formulate in a loss function. More details can be found in [39]. From their experiments, it has been shown that the model using the MSE along the two other losses performs better [39].

Additionally, as in [6], it has been investigated to add the beam configuration as input of the Pareto model proposed by Dan Nguyen et al. [11]. This allows for variable beam angles and a variable number of beams [40]. A model including the beam configuration in its inputs is performing better for variable beam configurations [40].

Furthermore, alternative model architectures to the HD U-Net are also working well in predicting the Pareto dose distribution. As an example, the standard U-Net has been used in [39, 40]. In [41], a Residual Network has been tested. Instead of predicting one Pareto optimal dose from scratch, it predicts an update of the initialized dose. This was done to remove nonlinearities from the dose prediction problem and reduce the difficulty of the dose prediction problem for the network. The general knowledge of dose distribution is used to initialize the dose given as input of the Residual Network. The error of the dose prediction of the model for the prostate VMAT case seems to be reasonable [41]. For the training of this model, only 25 Pareto plans per patient were used accounting for trade-offs between target coverage and sparing of bladder and rectum [41].

In addition to the model developed by Dan Nguyen et al. [11], another Pareto framework has been investigated by Tianfang Zhang et al. in [12]. It uses dose statistics to produce a tilting of the DVH of the neutral model predicting a single optimal dose. The U-Net model architecture is used along with the MSE, the DVH loss and the gradient loss. The gradient loss accounts for a smooth variation of the dose distribution. The neutral model is first trained for 100 epochs and then retrained with a modified DVH loss for 10 epochs.

The modified DVH loss is obtained by doing a DVH loss compared to a tilted ground truth DVH curve. This retraining step allows the model to deviate more aggressively toward one particular trade-off. The predicted dose resulting from the tilted model presents a better OAR sparing. However, to have a good accuracy, they used a separate model for each OAR sacrificing the inter-dependency of the ROIs [12]. Moreover, the level of dose reduction in the OAR in the predicted dose distribution cannot be set in the inputs of the tilted model [12].

Finally, a last model has been presented in [13] as an alternative to the model presented by Dan Nguyen et al. [11]. This model also takes the PTV/OAR trade-offs as input of the model to predict a Pareto optimal dose distribution. But, instead of using weights as input, it uses the desired DVH curves as input of the model. The weights are abstract and difficult to quantify. They are related to dose volume metrics but are not equal to them. This makes the model less accessible to manipulate by physicians. A solution is to allow the physicians to fine-tune the desired DVH curve. The proposed workflow is shown in Figure 3.5. First, the conventional model, presented in Section 3.2.1, predicts a single optimal dose distribution. If the corresponding DVH curves are not as desired, the physician can fine-tune the predicted DVH. Then, the refined DVH is used as an input of the proposed model and a new dose is predicted. In case the corresponding predicted DVH curves are still not as desired, the physician tunes them again and so on. A limitation of this model is that the tuned DVH might be totally unrealistic. The model then will fail to predict a Pareto optimal dose distribution [13].

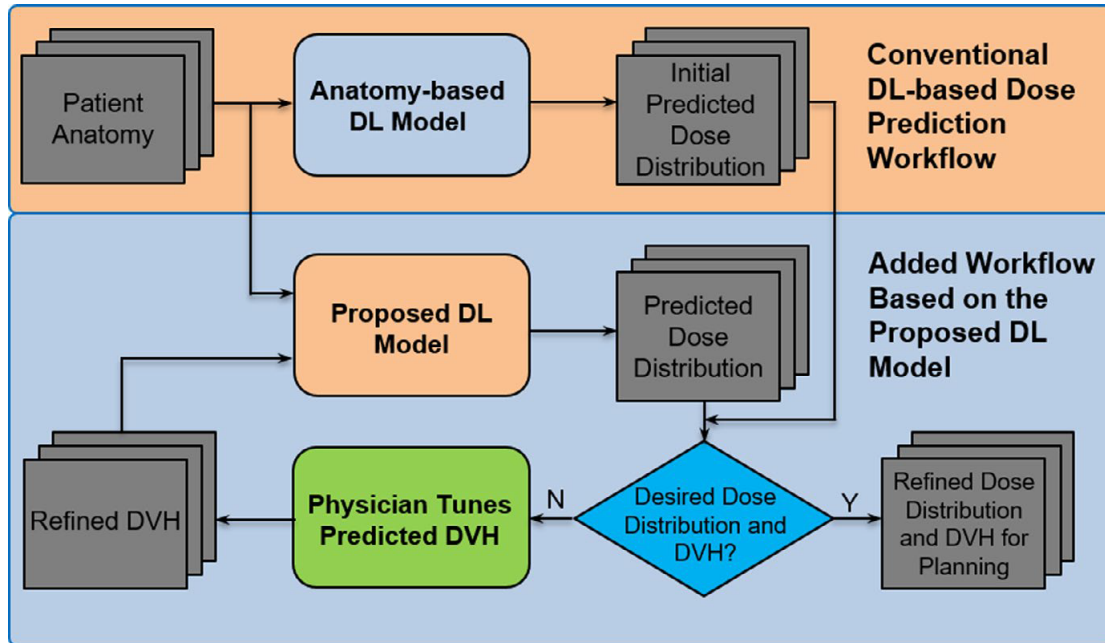


Figure 3.5: Workflow of Pareto dose prediction using the DVH curves as a quality assessment, allowing the physician to fine-tune them and run a new model again with the tuned DVH as an input. Credits: [13]

4 | Methods

This chapter presents the different methods used in this thesis. First, the patients database is presented. Then, the method used to generate artificial Pareto dose distributions is explained. Then, the architecture of the network is presented followed by the presentation of the different dose prediction models used. Finally, the method used for dose mimicking is presented followed by the formula for computing the Normal Tissue Complication Probability (NTCP) for xerostomia.

4.1 Patients Database

We have an anonymized dataset of 60 patients with oropharyngeal cancer who have been treated by IMPT for a fixed beam configuration. Four beams were used with the following beam configurations: $(10^\circ, 60^\circ)$, $(10^\circ, 120^\circ)$, $(350^\circ, 240^\circ)$ and $(350^\circ, 300^\circ)$ for the couch and gantry angles, respectively. The CTV has two dose prescription levels of 54.25 Gy and 70 Gy. The IMPT treatment plans were generated in RayStation (RaySearch Laboratories, Sweden) using a Monte Carlo dose calculation and a robust worst-case optimization on the CTV with 21 scenarios accounting for a 2.6 % proton range error and a 4 mm systematic setup error.

There are 13 OARs for these head and neck cancer patients: the brainstem, the upper esophagus, the glottic area, the oral cavity, the left parotid, the right parotid, the inferior pharyngeal constrictor muscle, the middle pharyngeal constrictor muscle, the superior pharyngeal constrictor muscle, the spinal cord, the left submandibular gland, the right submandibular gland and the supraglottic larynx.

For each patient, we have the following NumPy array files:

- The size/shape of the patient's CT: `arrayshape.npy`
- The CT scan: `ct.npy`
- A table saying which organs the patient has: `contoursexist.npy`. For example, 2 patients in the dataset do not have a left submandibular gland and one does not have a right submandibular gland.

- The mask of the CTV with the two prescription levels: `ctv.npy`
- The 3D mask of all OAR and body together: `struct.npy`
- The clinically delivered 3D dose: `dose.npy`

Additionally, the files `sample_probability_col.npy`, `sample_probability_row.npy` and `sample_probability_slc.npy` give the probability of having the CTV in a given row, column or slice for each patient. These files are used when using patching described in Section 4.3. The clinically delivered dose, the computed tomography and all contour masks have a voxel resolution of $3 \times 3 \times 3 \text{ mm}^3$.

In this thesis, the focus is set on lowering the dose in the left parotid, the right parotid, the left submandibular gland, and/or the right submandibular gland because these four organs are salivary glands involved in the xerostomia disease. The average volumes of these four organs, compared to the volume of the CTV, are given as a percentage in Table 4.1. The table shows that the volumes of these four organs are, on average, 7 to 27 times smaller than the CTV volume over all patients in the dataset.

	Left parotid	Right parotid	Left submandibular gland	Right submandibular gland
Volume compared to the CTV volume $\mu \pm \sigma$ [%]	$13.3 \pm 4.9 \%$	$13.4 \pm 4.5 \%$	$3.7 \pm 1.5 \%$	$3.9 \pm 1.4 \%$

Table 4.1: Average volumes of the left parotid, right parotid, left submandibular gland and right submandibular gland as a percentage of the two CTV volumes together.

4.2 Artificial Pareto Dose Generation

Generating a MCO database for each patient is a time-consuming process. Using a treatment planning system such as RayStation (RaySearch Laboratories, Sweden) would take up to 30 minutes to generate one plan, especially in the case of proton therapy, where robust optimization is slower than conventional (PTV-based) optimization (because we need to evaluate the dose in all scenarios). Reaching 1200 plans per patient as in [11] was not realizable with the current software available in the MIRO laboratory.

As an alternative, we have decided to generate artificial Pareto dose distributions for each patient by reducing the dose level in one organ at a time with a certain scaling. Even if it is not a clinically deliverable dose, we assumed it could give good results. The computation of one scaled dose distribution takes up to 2 seconds. The method is the following:

1. Take the boolean mask of the organ in which we want to reduce the dose
2. Perform a dilation of the mask by 2 voxels (6 mm²) in the three space dimensions (x, y, z)
3. Perform an inversion of the mask, replacing False values by 1 and replacing True values by the scaling factor
4. Apply a Gaussian filter with a standard deviation of 2 voxels to make sure there is no steep gradient of dose
5. Replace all values out of the body by 1
6. Multiply the ground truth dose distribution by the obtained mask

The dilation of the mask in step 2 is there to make sure the smoothing of the scaling takes place around the organ and not inside of it. Step 5 is there to make sure the dose distribution around the body stays the same compared to the ground truth dose distribution. The obtained results for a scaling of 0.7 in the right parotid for a given patient can be observed in Figure 4.1. The corresponding DVH is shown in Figure 4.2. There is not only a shift of the DVH curve for the right parotid but also for the two CTV levels.

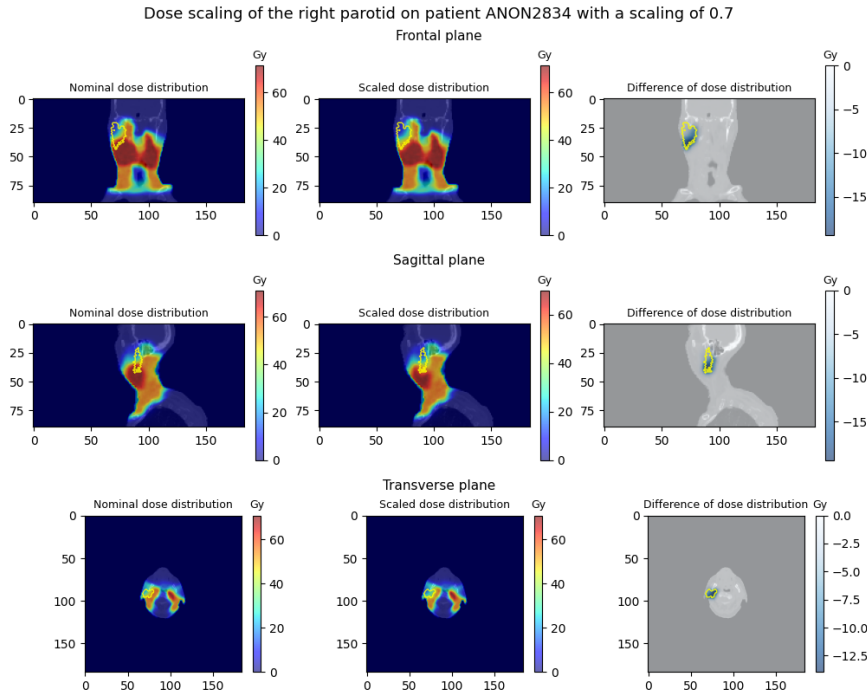


Figure 4.1: Dose distribution of the ground truth dose (left), of the 0.7 scaled dose in the right parotid (middle) and the difference of the two doses (right) for patient ANON2834. The right parotid is contoured in yellow in the images.

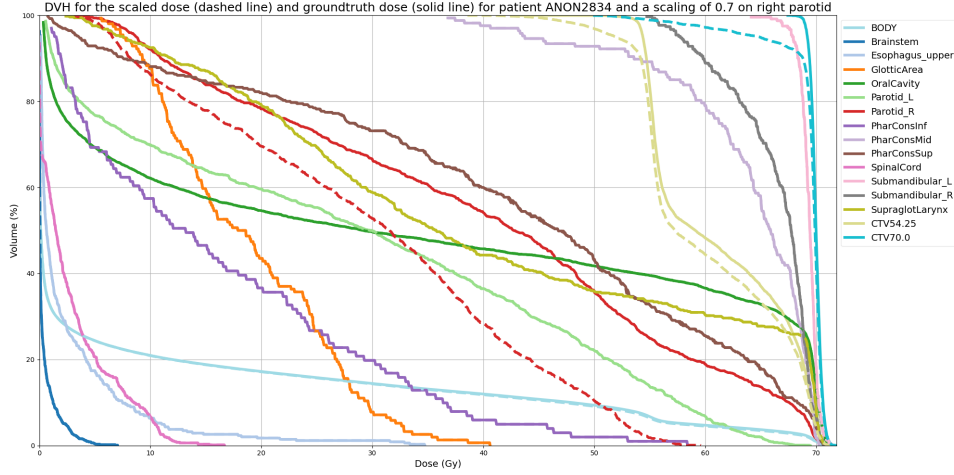


Figure 4.2: DVH for the scaled dose in the right parotid of 70% (dashed line) compared to the ground truth dose (solid line) for patient ANON2834.

4.3 Network Architecture

We work with a 5 layers HD U-Net architecture as presented in Figure 4.3 and described in Section 3.2.1. To update the different weights of the network, the Adam optimization algorithm from Keras is used. It is a stochastic gradient descent method that adaptively estimates the first-order moments and second-order moments by adapting the learning rate [42]. The Adam optimizer is used during the training to find the combination of weights that minimize the Mean Squared Error (MSE) between the predicted dose and the ground truth dose. For N voxels, the MSE is given by Equation 4.1.

$$MSE(y_{pred}, y_{true}) = \frac{1}{N} \sum_{i=1}^N (y_{pred}^i - y_{true}^i)^2 \quad (4.1)$$

The dense convolution uses the ReLU as an activation function and a zero padding to have a feature map of the same size as the input. The kernel size is (3,3,3) and each dense convolution operation contains 16 filters. The downsampling operations are max pooling operations with a (2,2,2) filter and a stride of (2,2,2) to reduce the feature maps by half of each of their dimensions.

As inputs of the architecture, we have 16 channels for each patient:

- The computed tomography
- The mask of the two CTVs together with their prescription levels
- The boolean mask of the body and the boolean masks for the 13 OARs

The output of the architecture model is the predicted 3D dose distribution in Gray.

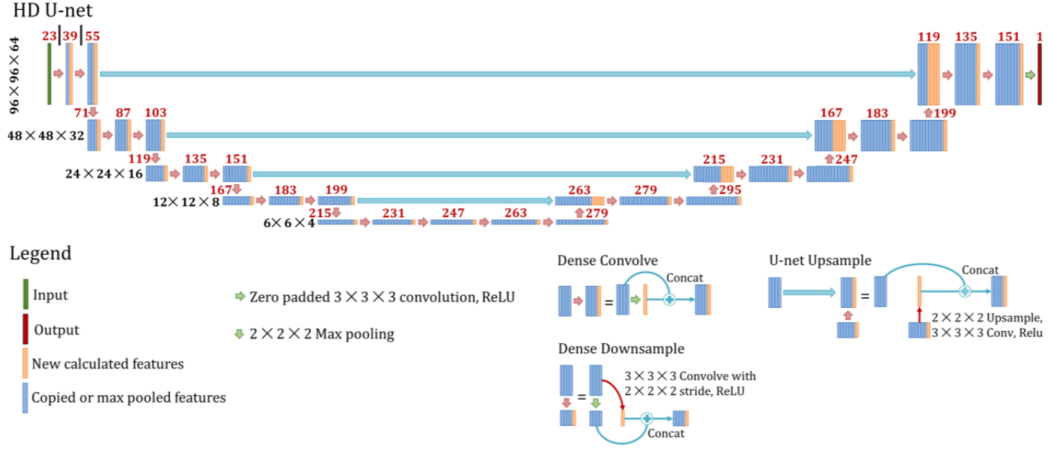


Figure 4.3: Hierchically Dense U-Net. The red numbers above the boxes represent the number of feature maps and the black numbers on the left represent the dimensions of each feature map. Credits: [7]

To speed up the training of the model and reduce overfitting, a patching system is used. Instead of feeding the model with full images, only a part of them is set as input to the model. At each iteration, a new randomly chosen patch of size $96 \times 96 \times 64$ is applied to the input channels. Using the `sample_probability` files, we make sure the patch contains the CTV. At each iteration, a new patch is computed reducing the chances of the model seeing twice the same images for a given patient. The model, therefore, predicts a patch containing only a part of the dose distribution. The complete dose distribution must be reconstructed from different predicted samples. The samples have a patch size of $96 \times 96 \times 64$ as the patched input images and are chosen with a stride of $32 \times 32 \times 32$ to cover the complete image. When there is an overlap between two predicted patches, the average between the overlapping predictions is taken.

4.4 Dose prediction models

Four different types of deep learning dose prediction models are presented in this thesis. First, a conventional dose prediction model is presented to benchmark our results. Secondly, a scaled dose prediction model is presented. Then, a model aiming at predicting different levels of scaling of one organ is used. Finally, a model allowing to scale one of the four salivary glands with a single model is presented. The first three models are trained for 100 epochs and the last one is trained for 250 epochs. All models are trained on 40 patients, validated with 10 patients, and tested with 10 patients.

4.4.1 Conventional dose prediction

A conventional dose prediction model (our baseline model) is necessary to be able to compare the results of the proposed Pareto dose prediction models. This model predicts a single optimal dose distribution as presented in Section 3.2.1.

To reduce overfitting and increase the size of the training database, data augmentation is used in the training step. This is possible due to the approximative horizontal symmetry of the head and neck anatomy. A flip of the patient is done for the training patients. The original patient data and the flipped data are set in the same batch resulting in a batch size of 2 and 40 training steps.

4.4.2 Scaled dose prediction

The scaled dose prediction model is useful when needing to predict a dose distribution in a certain organ with a lower dose level than that of the clinical ground truth in the database.

The model is trained with different (input, output) pairs from the conventional model. The desired output we use to train the model is the artificial Pareto dose presented in Section 4.2 with a scaling in the desired organ. We will call this desired output the scaled ground truth. Each model is organ-specific and will be able to lower the dose in a single organ. A batch size of 1 is used for the training.

In the results, two models are presented. The first model lowers the dose in the right parotid by a scaling to have 70% of the dose and the second model lowers the dose in the right submandibular gland by a scaling to have 70% of the dose.

4.4.3 Dose prediction with different scalings in a single organ

Due to the large amount of OARs around a head and neck cancer and their relatively small size, an organ-specific model is first proposed.

Our model aims at predicting the dose distribution with different scalings for a given organ. The desired scaling is set as an input of the model and the corresponding scaled dose is predicted. As an example, we take a model predicting different scaled doses in the right parotid.

This model is also trained with different (input, output) pairs than the conventional model. The desired scaling is set in the inputs by changing the value of the boolean mask of the right parotid from 1 to the scaling as proposed in [11]. The model is then trained using the artificial Pareto dose distribution, presented in Section 4.2, corresponding to the desired scaling.

The model is trained to predict a dose distribution with a scaling between 0.5 and 1.5 in the right parotid. In each batch, 5 different scenarios for the selected patient are given resulting in different (input, output) pairs. The batch size is thus 5. The first scenario is the nominal dose prediction with no scaling. The four other scenarios are generated with randomly chosen scalings between 0.5 and 1.5. The discretization of this interval is done by taking points equidistant to 0.05 from each other.

4.4.4 Pareto dose prediction

The aim of this thesis was to reach a Pareto dose prediction model for multiple organs as presented by Dan Nguyen et al. in [11]. Here, a last model is presented as a generalization of the organ-specific model presented in Section 4.4.3. Instead of having a model able to predict a scaled dose in one organ only, this model can perform a scaling of the dose in one of the four following organs: the left and right parotids and the left and right submandibular glands. These four organs are the organs involved in xerostomia

The model is trained to predict a dose distribution with a scaled dose in one of the four organs. The scaling is between 0.5 and 1.5 and is set as an input in the mask of the organ for which the scaling is desired. As proposed in [11], the ones in the mask of this organ are set to the desired scaling.

For the training, the batch size is 4. In each batch, we consider four different scenarios for the selected patient. Each scenario considers a randomly chosen scaling of one of the four organs. The randomly chosen scaling is chosen in the interval between 0.5 and 1.5 with equidistant numbers 0.05 from each other. Here, no scenario accounts for the nominal case without any scaling. We suppose the nominal case is obtained when the randomly chosen scaling is 1.

4.5 Dose mimicking

Since our simplistic model to generate scaled ground truth doses might lead to highly unrealistic dose distributions, a dose mimicking step has been applied to the predicted dose to evaluate its interest. This mimicking process has been done for the dose predicted with a lower dose in the right parotid and on the dose predicted with a lower dose intensity in the right submandibular gland for all test patients.

Dose mimicking is an automatic method that finds the machine parameters allowing us to get a clinically deliverable dose distribution as close as possible to the predicted dose. The beam configuration is the same for all patients.

The dose mimicking step is performed through a script in the RayStation software (RaySearch Laboratories, Sweden). Isodose volumes were generated from the predicted dose every 2 Gy within 95%–107% of the prescription dose and 5 Gy outside. An example of isodoses for a patient is presented in Figure 4.4.

Subsequently, an isodose-based optimization was performed using minimum and maximum objectives on the CTVs and on each isodose volume. The isodoses above 70% of the prescribed dose had maximum and minimum dose objectives while isodoses below 70% of the prescribed dose only had a maximum objective. This allows placing the spots in a restricted area around the CTV to prevent the optimizer to set spots in the whole patient volume.

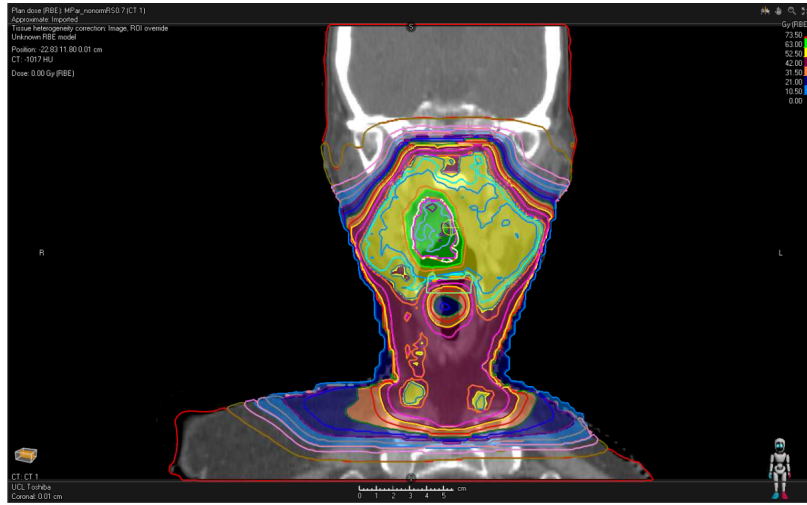


Figure 4.4: Example of isodose volumes for a certain dose distribution in the RayStation software.

Among the objectives, some of them were set as robust to account for treatment uncertainties. The robust optimization accounts for a proton range uncertainty of 2.6% and a setup error of the patient of 4 mm in the three space dimensions (Inferior-Superior, Right-Left, Anterior-Posterior). The following objectives are set as robust:

- On the CTV volumes: Minimum objectives at the prescribed doses: 70 Gy for the primary CTV and 54.25 Gy for the two nodal CTV volumes (CTVnR and CTVnL). Two maximum objectives on the primary CTV were set as robust in order to avoid hot spots: one on the CTV primary and another one on an expanded version of 3cm of the CTV.
- On the isodose volumes a total of four robust objectives were set, two minimum objectives per dose prescription level: 6950 and 5153 corresponding to the closest value to 100% of the prescribed dose to the CTV and 6650 and 5153 corresponding to 95% of the prescribed dose.

The spot weights were chosen empirically after manual tuning in a subset of the patient data set. The final parameters were saved as the template that was used later to run the

optimization (mimicking) for all test patients. The template containing all optimization objectives for the dose mimicking is presented in Figure 4.5.

The plan was optimized through an iterative process using Monte Carlo dose calculation.

Uniform dose	Beam set (RBE)	CTV_LPO	Uniform dose 0.01 Gy (RBE), Beam 'LPO'		0.00
Uniform dose	Beam set (RBE)	CTV_RPO	Uniform dose 0.01 Gy (RBE), Beam 'RPO'		0.00
Uniform dose	Beam set (RBE)	CTV_LAO	Uniform dose 0.01 Gy (RBE), Beam 'LAO'		0.00
Uniform dose	Beam set (RBE)	CTV_RAO	Uniform dose 0.01 Gy (RBE), Beam 'RAO'		0.00
Min dose	Plan (RBE)	CTVp_7000	Min dose 70.00 Gy (RBE)	★	400.00
Max dose	Plan (RBE)	CTVp_7000	Max dose 70.00 Gy (RBE)	★	400.00
Min dose	Plan (RBE)	CTVnL_5425	Min dose 54.25 Gy (RBE)	★	400.00
Min dose	Plan (RBE)	CTVnR_5425	Min dose 54.25 Gy (RBE)	★	400.00
Max dose	Plan (RBE)	CTV_7000_3cm	Max dose 70.00 Gy (RBE)	★	2000.00
Min dose	Plan (RBE)	CTVnR_7000	Min dose 70.00 Gy (RBE)	★	400.00
Max dose	Plan (RBE)	CTVnR_7000	Max dose 70.00 Gy (RBE)	★	400.00
Max dose	Plan (RBE)	isodoses 7350	Max dose 75.50 Gy (RBE)		100.00
Min dose	Plan (RBE)	isodoses 7150	Min dose 71.50 Gy (RBE)		50.00
Max dose	Plan (RBE)	isodoses 7150	Max dose 73.50 Gy (RBE)		10.00
Min dose	Plan (RBE)	isodoses 6950	Min dose 69.50 Gy (RBE)	★	100.00
Max dose	Plan (RBE)	isodoses 6950	Max dose 71.50 Gy (RBE)		50.00
Min dose	Plan (RBE)	isodoses 6750	Min dose 67.50 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 6750	Max dose 69.50 Gy (RBE)		10.00
Min dose	Plan (RBE)	isodoses 6650	Min dose 66.50 Gy (RBE)	★	100.00
Max dose	Plan (RBE)	isodoses 6650	Max dose 67.50 Gy (RBE)		1.00
Min dose	Plan (RBE)	isodoses 6150	Min dose 61.50 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 6150	Max dose 66.50 Gy (RBE)		10.00
Min dose	Plan (RBE)	isodoses 5696	Min dose 56.96 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 5696	Max dose 61.50 Gy (RBE)		10.00
Min dose	Plan (RBE)	isodoses 5496	Min dose 54.96 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 5496	Max dose 56.96 Gy (RBE)		10.00
Min dose	Plan (RBE)	isodoses 5296	Min dose 52.96 Gy (RBE)	★	100.00
Max dose	Plan (RBE)	isodoses 5296	Max dose 54.96 Gy (RBE)		50.00
Min dose	Plan (RBE)	isodoses 5153	Min dose 51.53 Gy (RBE)	★	100.00
Max dose	Plan (RBE)	isodoses 5153	Max dose 52.96 Gy (RBE)		20.00
Max dose	Plan (RBE)	isodoses 4653	Max dose 51.53 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 4153	Max dose 46.53 Gy (RBE)		50.00
Max dose	Plan (RBE)	isodoses 3653	Max dose 41.53 Gy (RBE)		50.00
Max dose	Plan (RBE)	isodoses 3153	Max dose 36.53 Gy (RBE)		50.00
Max dose	Plan (RBE)	isodoses 2653	Max dose 31.53 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 2153	Max dose 26.53 Gy (RBE)		10.00
Max dose	Plan (RBE)	isodoses 1653	Max dose 21.53 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 1153	Max dose 16.53 Gy (RBE)		5.00
Max dose	Plan (RBE)	isodoses 653	Max dose 11.53 Gy (RBE)		1.00
Max dose	Plan (RBE)	isodoses 153	Max dose 6.53 Gy (RBE)		5.00
Max dose	Plan (RBE)	isodoses 0	Max dose 1.50 Gy (RBE)		1.00

Figure 4.5: Screenshot from RayStation showing the template for the different objectives used in the optimization for the dose mimicking. The first column gives the type of objective, the third gives the isodose name, the fourth gives the detailed objective, the fifth has a star if the objective is robust and the last column gives the spot weight used for each objective.

4.6 Normal Tissue Complication Probability

To evaluate the mimicked plans for sparing the salivary glands, the Normal Tissue Complication Probability (NTCP) for xerostomia is computed. We compare the NTCP mimicked plans for a lower dose in the right parotid and a lower dose in the submandibular gland with the NTCP for the clinical ground truth. The NTCP is a percentage. Different grades of side effects exist and we consider grades 2 and 3. The formula for computing the NTCP is given by Equations 4.2 and 4.3 from [43]. The required parameters are presented in Table 4.2. $Dmean_{RP}$ is the mean dose value in the right parotid, $Dmean_{LP}$ is the mean dose value in the left parotid, $Dmean_{RS}$ is the mean dose value in the right submandibular gland and $Dmean_{LS}$ is the mean dose value in the left submandibular gland. V_{LS} and V_{RS} are the volumes of the left and right submandibular gland, respectively.

$$NTCP(S) = \frac{1}{1 + e^{-S}} \quad (4.2)$$

with

$$S = \beta_0 + \bar{\beta} + \beta_{Parotid} \cdot (\sqrt{Dmean_{RP}} + \sqrt{Dmean_{LP}}) + \beta_{Submandibular} \cdot \frac{Dmean_{LS} \cdot V_{LS} + Dmean_{RS} \cdot V_{RS}}{V_{LS} + V_{RS}} \quad (4.3)$$

Parameter	Grade 2	Grade 3
β_0	-2.2951	-3.7286
$\overline{\beta}$	0	0
$\beta_{Parotid}$	0.0996	0.0855
$\beta_{Submandibular}$	0.0182	0.0156

Table 4.2: Parameters to compute the NTCP

5 | Results

This chapter presents the different results obtained in this thesis. First, the analysis of the dose predicted by the four dose prediction models has been done on the test patients. The first model is the conventional dose prediction model presented to benchmark our results. The second model is the model that allows reducing the dose in an organ. In the results, we analyze the cases where the right parotid gland or the right submandibular gland have a dose decreased by 30%. The third model is a model that allows to change the dose in the right parotid by a scaling set in the inputs of the model. The last model is a model allowing to change the dose in one organ among the four salivary glands by a scaling between 0.5 and 1.5. Then, the results of the dose mimicking are presented. The dose mimicking was done on the second model allowing either to decrease the dose in the right parotid gland or to decrease the dose in the right submandibular gland. Finally, the Normal Tissue Complication Probability (NTCP) for xerostomia is computed for the mimicked plans and compared to the reference plan.

To evaluate the different dose distributions over the whole set of patients, boxplots are used. A boxplot is a graph that enables to visualize the statistics over a dataset. The box contains the middle 50% of the data and is covering the interquartile range. The median is the line drawn inside this box. The whiskers extend the box to 1.5 of the interquartile range. The dots further away from the box and whiskers are outliers.

5.1 Dose prediction models

5.1.1 Conventional dose prediction

The conventional model predictions have an average MSE of 1.03 over the 10 test patients. The model takes around 7 hours to be trained for 100 epochs.

Figure 5.1 presents the boxplots for the error of the mean dose for the OARs and the error of the D_{95} for the CTVs. From this figure, we observe that the interquartile range lies between -5 Gy and +5Gy for all ROIs and the medians are relatively close to 0. However, there is a lower target coverage than needed for the CTV7000.

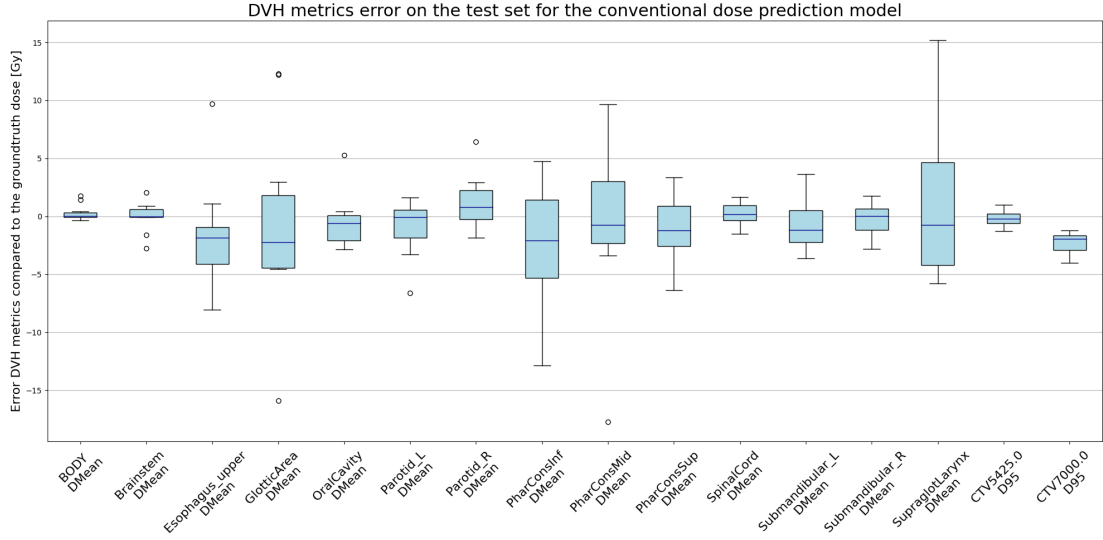


Figure 5.1: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution and the ground truth.

The best dose prediction is for patient ANON1169 with an MSE of 0.56. The predicted dose distribution for this patient is compared with the ground truth dose distribution in Figure 5.2. The difference in dose intensity ranges from -20 to +20 Gy. The biggest differences are found in the borders of the dose distribution and in regions with steep dose gradients (e.g. in the border of the CTV).

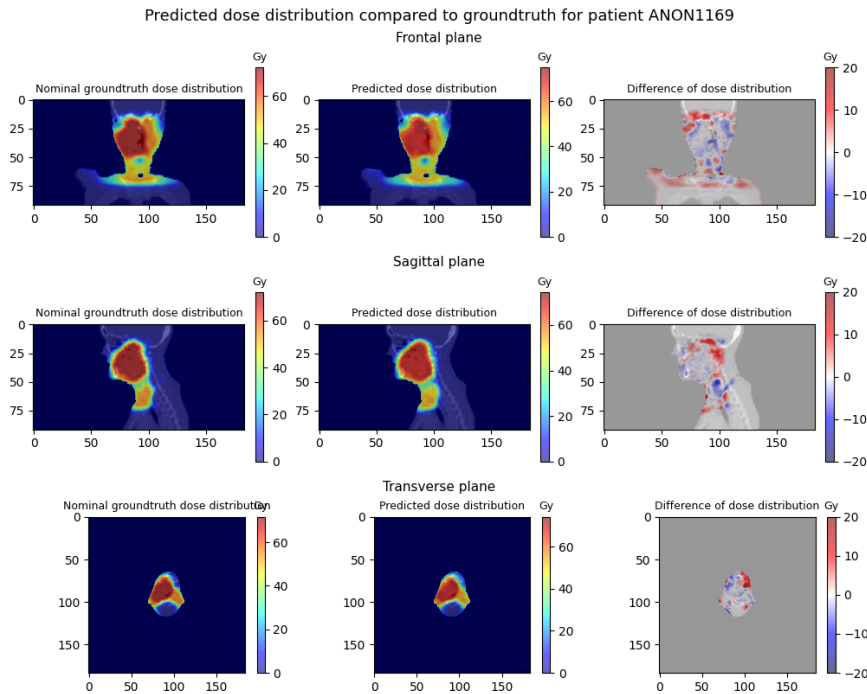


Figure 5.2: Comparison of the predicted dose distribution with the ground truth dose distribution for patient ANON1169.

5.1.2 Scaled dose prediction

5.1.2.1 Model predicting a lower dose in the right parotid

The model predicting a lower dose with a scaling of 0.7 in the right parotid, takes around 7 hours to be trained for 100 epochs. The average MSE for the predicted dose on the test set is 1.15.

Figure 5.3 presents the boxplots for the error between the predicted dose distribution and the nominal (blue) and scaled (green) ground truth for all ROIs. The D_{mean} is used as a comparison DVH metric for the OARs and the D_{95} is used to evaluate the target coverage of the CTVs. Overall, the error seems to be of the same order of magnitude as for the conventional model. We observe that the right parotid has a lower mean dose with a median of 5 Gy compared to the nominal ground truth and has a small mean dose error compared to the scaled ground truth.

The best dose prediction is for patient ANON1089 with an MSE of 0.62. Figure 5.4 shows the comparison between the predicted dose distribution for this patient and the nominal ground truth dose distribution. There are more pronounced spots with a lower dose than for the conventional model. In particular, we observe that the dose has been reduced by 10 to 20 Gy in and around the right parotid (contoured in yellow on the images).

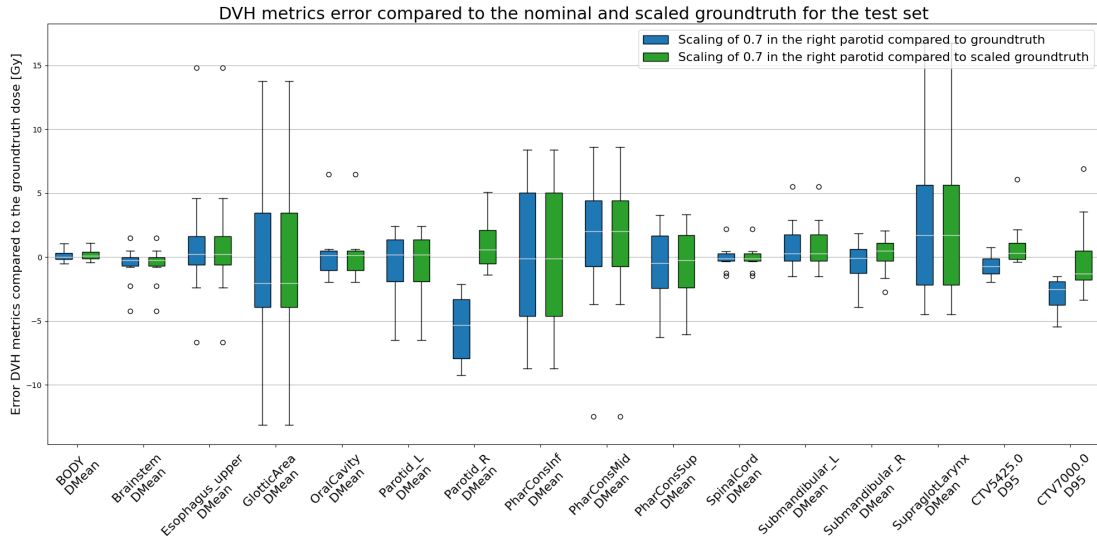


Figure 5.3: Boxplots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution and the nominal ground truth (blue) or the scaled ground truth with a lower dose in the right parotid (green).

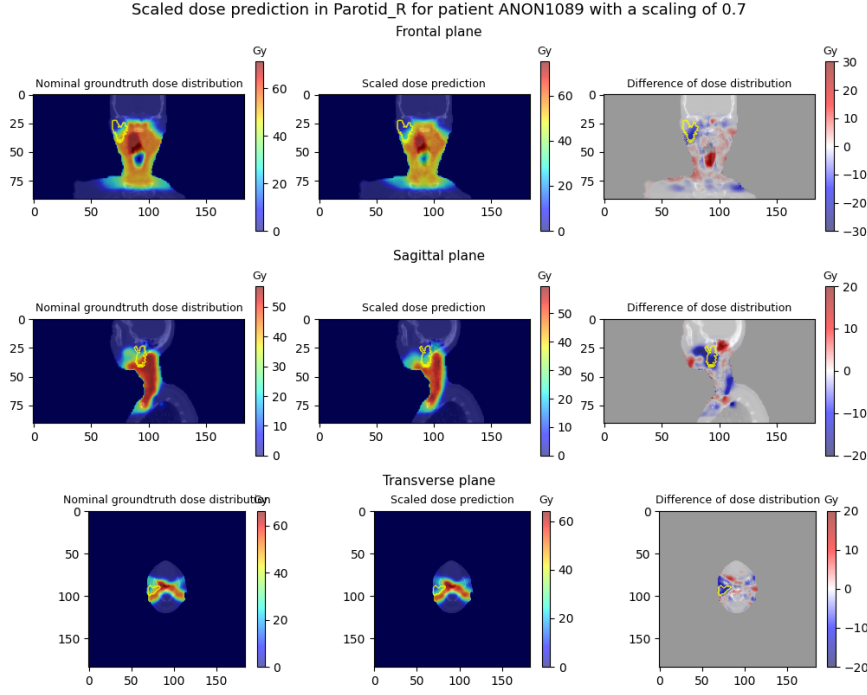


Figure 5.4: Comparison of the predicted dose distribution having a lower dose in the right parotid (middle) with the nominal ground truth dose distribution (left) for patient ANON1089. The right parotid is contoured in yellow on the images.

5.1.2.2 Model predicting a lower dose in the right submandibular gland

The model that only predicts a scaled dose in the right submandibular gland takes around 7 hours to be trained as well. The average MSE for the predicted dose on all test patients is 1.34.

Figure 5.5 presents the boxplots for the error on the D_{mean} and D_{95} for the OARs and CTVs respectively. It compares the predicted dose distribution for a lower dose in the right submandibular gland with the nominal (blue) and scaled (green) ground truths on all test patients. Overall, the error is the same when compared to the nominal or scaled ground truth. Two noticeable variations are observed. First, the predicted dose has a mean dose of almost 15 Gy less than the nominal ground truth and a very small error in the right submandibular gland compared to the scaled ground truth. The model is thus managing to lower well the dose in the right submandibular gland as desired. Secondly, there is a lowering of the dose in the CTV7000 compared to the nominal ground truth and not compared to the scaled ground truth. Due to the proximity of the submandibular gland with the CTV, lowering the dose in the right submandibular unavoidably reduces the minimal dose in the CTV. This happens for the prediction but also for the scaled ground truth.

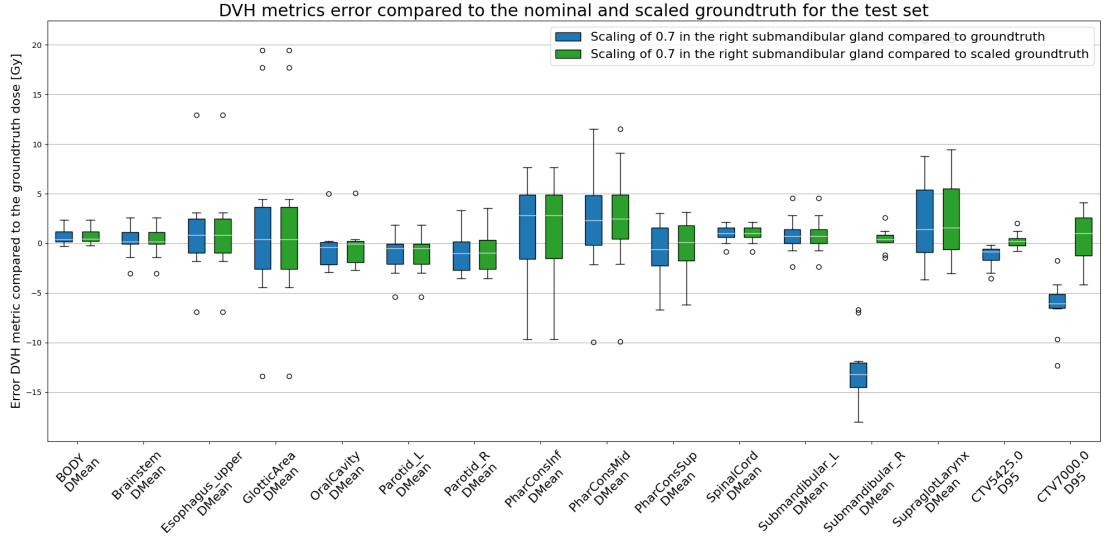


Figure 5.5: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution and the nominal ground truth (blue) or the scaled ground truth with a scaling of 0.7 in the right submandibular gland (green).

The best MSE of the predicted dose distribution compared with the scaled ground truth is obtained for the patient ANON1089. The MSE for this patient is 0.68. The comparison of the predicted dose distribution with the nominal ground truth dose distribution is presented in Figure 5.6. The dose is mostly lowered by up to 20 Gy in and around the submandibular gland.

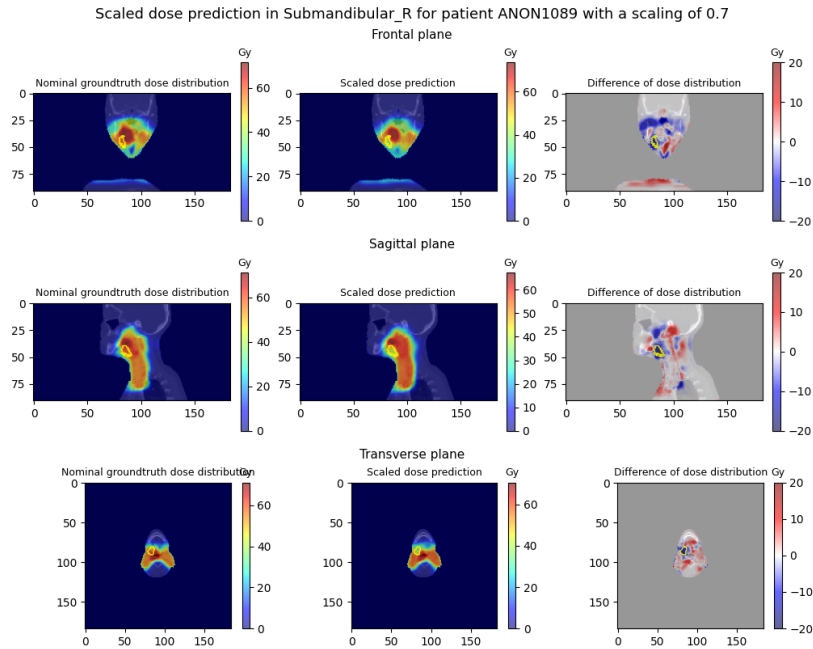


Figure 5.6: Comparison of the predicted dose distribution having a lower dose in the right submandibular gland with the nominal ground truth dose distribution for patient ANON1089. The right submandibular gland is contoured in yellow on the images.

5.1.3 Dose prediction with different scalings in a single organ

The model predicting a dose with different scalings in the right parotid takes around 2 days to be trained for 100 epochs. This is due to the fact that the scaled dose distribution for the desired output are scaled at each iteration with the randomly chosen scaling and that we have a batch size of 5 instead of a batch size of 2 for the conventional model. The average MSE over all test patients for the scalings [0.6,0.7,0.8,1,1.2,1.4] is 2.62 which is higher than for the previous models.

Figure 5.7 presents the boxplots for the DVH metrics error on the test set for the model dose predictions compared to the nominal ground truth. The considered metrics are the D_{mean} and D_{95} for the OARs and the CTVs, respectively. The desired scalings in the right parotid are 0.6, 0.7, 0.8, no scaling, 1.2 and 1.4. We observe that there is indeed a mean dose variation in the right parotid along with the variation of scalings. The scaling of 0.6 implies that the mean dose in the right parotid is lowered with a median of 6 Gy and the scaling of 1.4 implies that the mean dose in the right parotid is increased by 5 Gy compared to the reference.

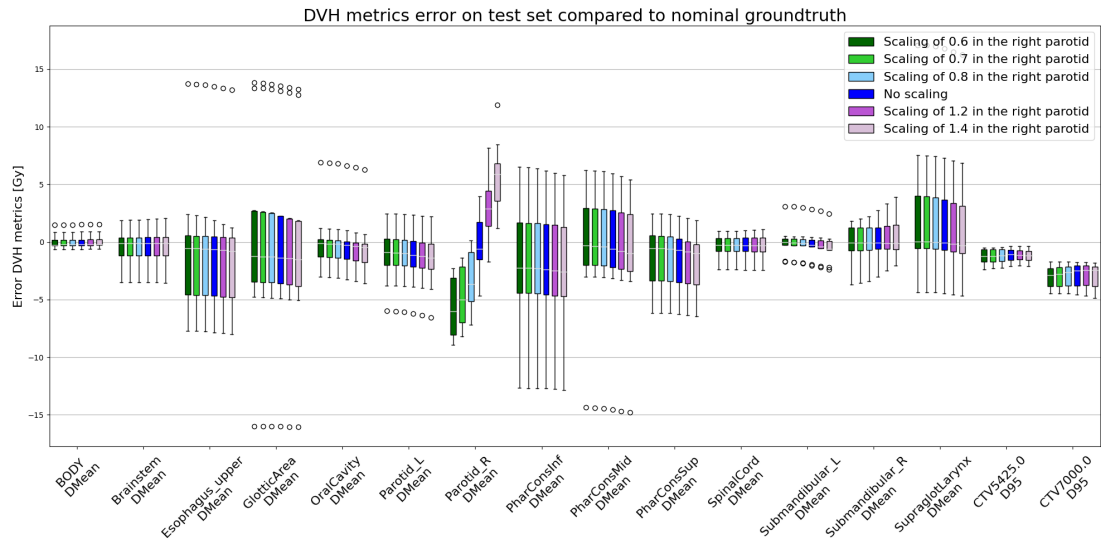


Figure 5.7: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution for different scalings of dose in the right parotid and the nominal ground truth.

Figure 5.8 presents the boxplots for the DVH metrics error on the test set for the model with different scalings in input compared to the corresponding scaled ground truth. The considered metrics are D_{mean} and D_{95} for the OARs and the CTVs, respectively. The desired scalings in the right parotid are 0.6, 0.7, 0.8, no scaling, 1.2 and 1.4. It can be observed that the mean dose error in the right parotid is quite consequent compared to the desired dose distribution. The mean dose in the right parotid is up to 10 Gy higher

than desired for a scaling of 0.6 in the right parotid and is up to 10 Gy lower for a desired scaling of 1.4 in the right parotid. The desired effect is thus attenuated in reality.

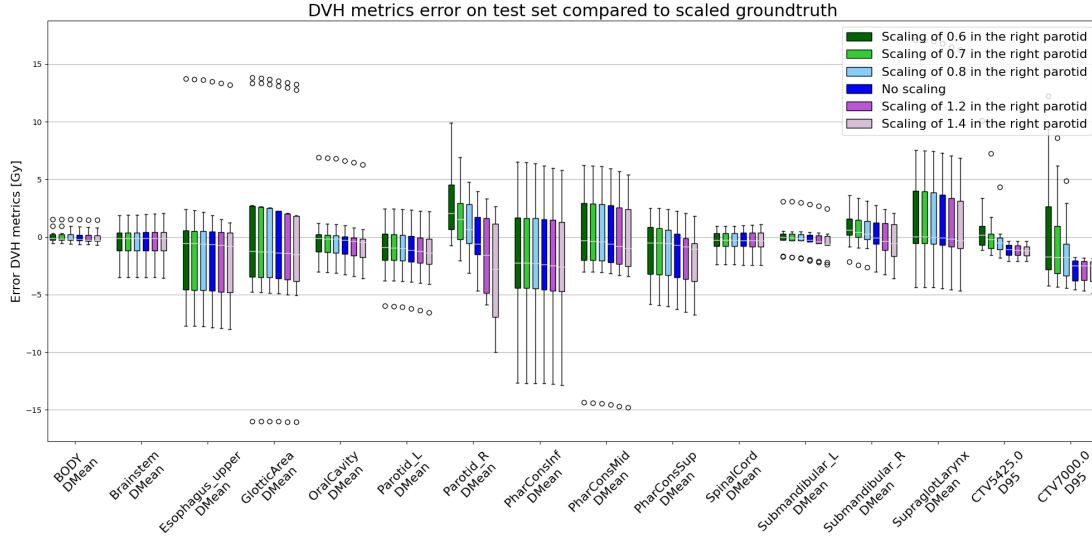


Figure 5.8: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution for different scalings of dose in the right parotid and the corresponding scaled ground truth.

Figure 5.9 presents the DVH curves for the right parotid for the patient ANON2146. It shows the different DVH curves corresponding to a scaling of 0.6, 0.8, 1, 1.2 and 1.4 in the right parotid. The desired DVH curve (solid line) and the predicted DVH curve (dashed line) for the right parotid are presented. For a scaling close to 1 (i.e. 0.8, 1 and 1.2), the prediction curve is very close to the one of the scaled ground truth. However, for a scaling of 0.6, we observe that the prediction curve is not shifted as much to the left as desired resulting in a higher dose in the right parotid.

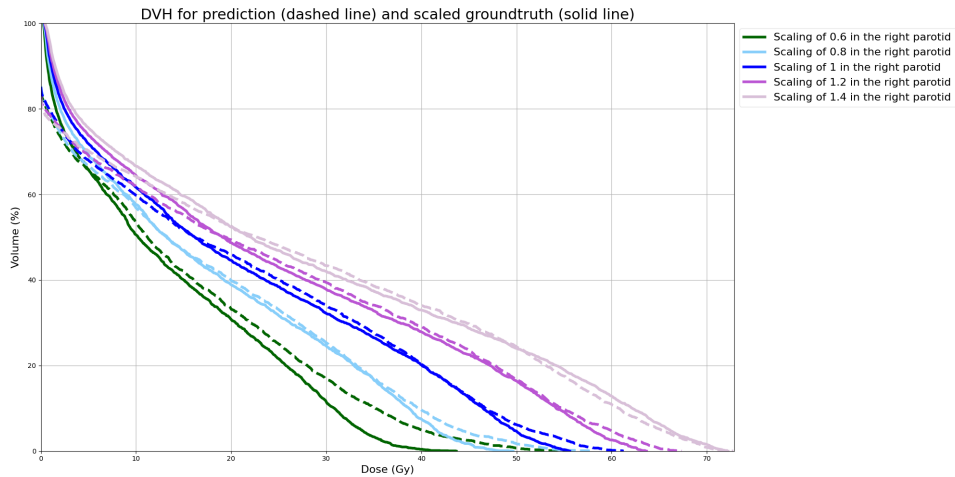


Figure 5.9: DVH for the prediction of a scaled dose in the right parotid (dashed line) compared to the scaled ground truth dose (solid line) for patient ANON2146.

5.1.4 Pareto dose prediction

The model predicting a scaled dose distribution in one of the four salivary glands with a scaling between 0.5 and 1.5 takes 2 days and 15 hours to be trained for 250 epochs. The average MSE on the test set for the predicted doses with a lower dose in each of the salivary glands with a scaling of $[0.6, 0.7, 0.8, 1, 1.2, 1.4]$ is 3.33.

Figure 5.10 presents the boxplots for some DVH metrics compared to the reference ground truth. The considered DVH metrics are the mean dose and the D_{95} for the OARs and CTVs, respectively. Only some of the OARs are presented, the others having the same errors as in previous models. For each of the four glands, the results are presented for a scaling of 0.6 and 1.4. These two scalings were not very working well in the organ-specific model for which the results are presented in Section 5.1.3.

Figure 5.11 presents the boxplots for some DVH metrics on the predicted scaled dose in a given organ compared to the corresponding scaled ground truth. The considered DVH metrics are the mean dose and the D_{95} for the OARs and CTVs, respectively. Only some of the OARs are presented, the others having the same errors as in previous models. For each of the four glands, the results are presented for a scaling of 0.6 and 1.4. These two scalings were not working very well in the organ-specific model for which the results are presented in Section 5.1.3.

Compared to the reference ground truth, we observe an increase or decrease of around 5 Gy for a scaling of 0.6 and 1.4 respectively in the desired organ. However, when no scaling is desired in a specific organ, the median seems to be slightly higher by 2 Gy than the reference.

When comparing the mean doses of the scaled dose in a particular gland to the corresponding scaled ground truth, we observe the error goes up to -20 and +20 Gy for a scaling of 0.6 and 1.4 respectively in the right submandibular gland. A similar slightly lower error appears when scaling the dose in the left submandibular gland. We can also observe that the model understands that a certain dose level has to be reached in the CTVs. But this implies that the desired scaling in the right submandibular cannot be reached and that there is a higher error in the CTV level compared to the scaled ground truth.

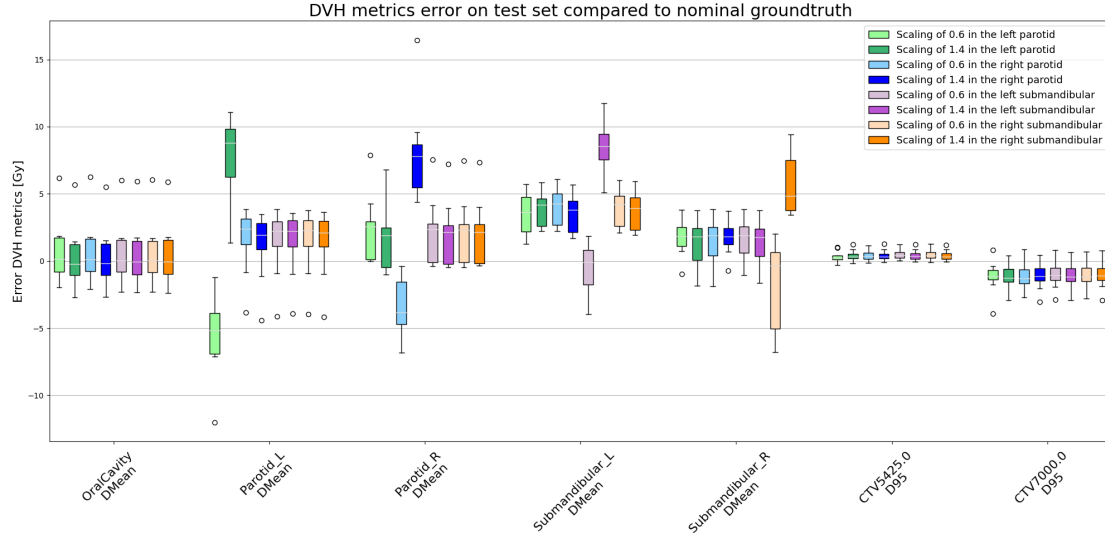


Figure 5.10: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution for different scalings of dose in the four glands and the reference ground truth.

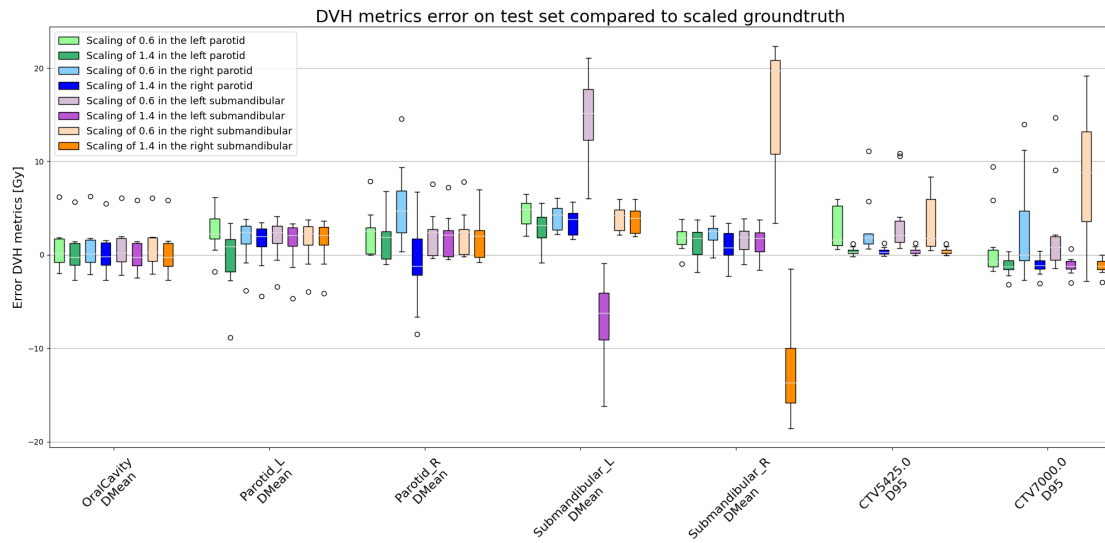


Figure 5.11: Boxpots for the error of the mean dose for OARs and D_{95} for the CTVs between the predicted dose distribution for different scalings of dose in the four glands and the corresponding scaled ground truth.

5.2 Dose mimicking

The dose mimicking step has been done on the predictions of the dose distribution from the models presented in Section 4.4.2 with the corresponding results in Section 5.1.2. The dose mimicking for one plan takes around 10 to 20 minutes.

Figure 5.12 presents the boxplots for the mimicked dose aiming at lowering the dose in the right parotid on all test patients. The D_{mean} and D_{95} are used as comparison metrics for the OARs and CTVs, respectively. The mimicked dose distribution is compared to the reference ground truth dose distribution (blue) and to the predicted dose distribution with a lower dose in the right parotid (green). Compared to the reference, the mimicked plan indeed has a lower dose of 2 to 5 Gy in the interquartile for the right parotid. But, as a reaction, it increases the dose in the left and right submandibular glands by up to 3 Gy in the interquartile. The error of the dose in the CTV is very small compared with the reference plan. Compared to the prediction, the mimicked plan increases the dose in the right parotid again by up to 5Gy. In general, the mimicked plan has a higher dose level in almost all OARs and CTVs.

Figure 5.13 presents the boxplots for the mimicked plan aiming at lowering the dose in the right submandibular gland on all test patients. The D_{mean} and D_{95} are used as comparison metrics for the OARs and CTVs, respectively. The mimicked dose distribution is compared to the reference ground truth dose distribution (blue) and to the predicted dose distribution with a lower dose in the right submandibular gland (green). Compared to the reference plan, the mimicked plan indeed has a lower dose in the right submandibular gland with a median of around 5 Gy lower. The other OARs have an interquartile increase of dose up to 5 Gy. But, the CTV7000 has a lower dose level with a reduction of up to 2.5 Gy. Such a reduction is mostly due to the closeness between the submandibular gland and the CTV7000. Compared to the predicted plan, the mean dose in the right submandibular gland is increased with a median of 7.5 Gy. This also results in an increase of the minimum dose level in the CTVs. Concerning the other OARs, there is an increase of up to 2.5 Gy.

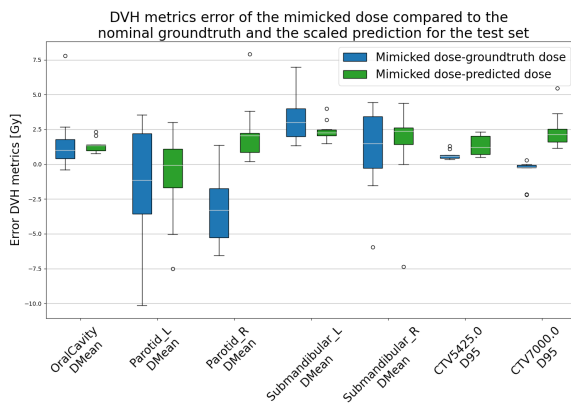


Figure 5.12: Boxplots of the DVH metrics error of the mimicked plan compared to the reference and the prediction when lowering the dose in the right parotid.

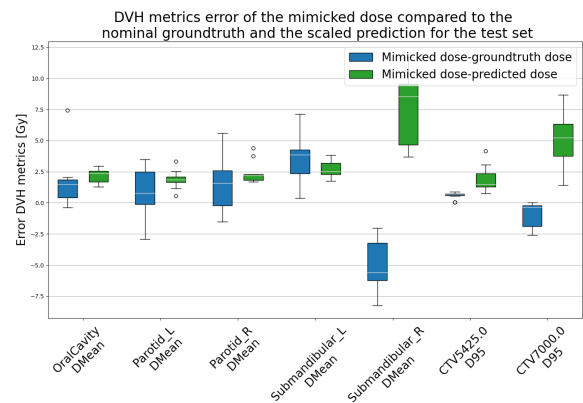


Figure 5.13: Boxplot of the DVH metrics error of the mimicked plan compared to the reference and the prediction when lowering the dose in the right submandibular gland.

As an illustration, the mimicked plans for the patient ANON1089 are analyzed further.

5.2.1 Mimicked plan for a lower dose in the right parotid for patient ANON1089

Figure 5.14 shows the dose distribution comparison between the mimicked plan (middle) and the predicted dose distribution for a lower dose in the right parotid (left) together with their difference in dose (right). The pink color accounts for at least 5 Gy more in the mimicked plan and the light blue for at least 5 Gy less. We observe many spots where the dose is increased by 5 Gy or more in the mimicked plan and particularly around the right parotid (indicated by the white arrow).

Figure 5.15 presents the dose distribution comparison between the mimicked plan (middle) and the reference ground truth dose distribution (left) together with their difference in dose (right). The pink color accounts for at least 5 Gy more in the mimicked plan and the light blue for at least 5 Gy less. We observe many spots where the dose is decreased by 5 Gy or more in the mimicked plan and particularly around the right parotid (indicated by the white arrow). However, due to the decrease of dose in the right parotid, there are some increases of 5 Gy or more elsewhere.

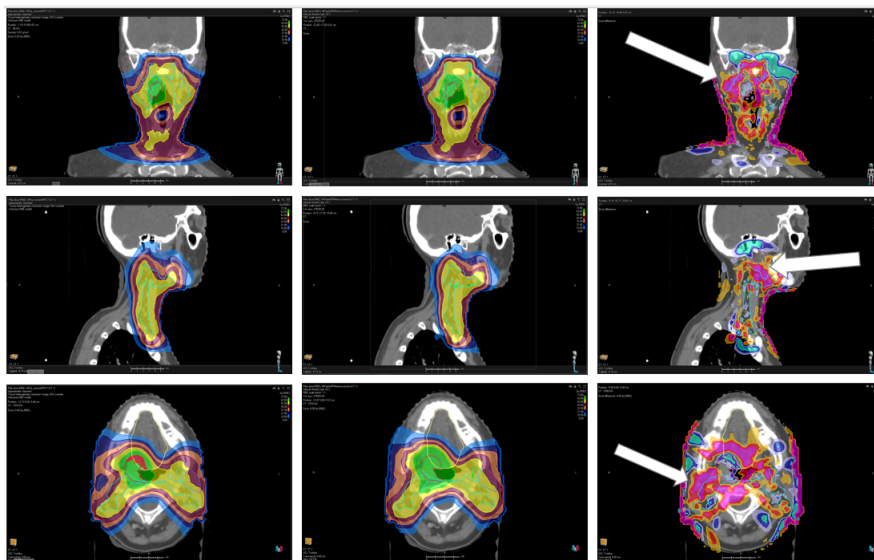


Figure 5.14: Comparison of the predicted dose distribution (left) with the mimicked plan lowering the dose in the right parotid (middle) and their difference in dose distribution (right) for ANON1089. The right parotid is indicated by white arrows.

Figure 5.16 presents the DVH for some ROIs with a comparison between the mimicked plan (solid line), the predicted dose distribution (dashed line) and the reference plan (dotted line). Some relevant dose volume metrics for the comparison of the three plans are

presented in Table 5.1.

For the mimicked plan, we observe that the dose level in the right parotid is significantly decreased with a D_{mean} decrease of 6 Gy but it has a bit increased compared to the predicted dose. The left parotid has a slightly lower dose in the mimicked plan compared to the two others as well. However, to compensate, the mean dose in the left submandibular gland is increased by 2 Gy in the mimicked plan compared to the two others. The dose in both submandibular gland seems to increase when looking at their DVH curves (in brown). The target coverage stays good in the mimicked plan and is even better for the CTV5425.

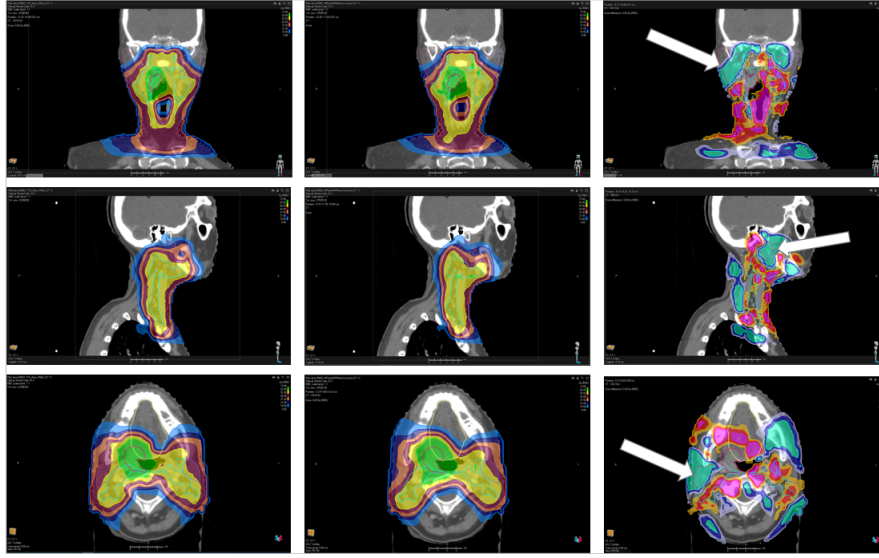


Figure 5.15: Comparison of the ground truth dose distribution (left) with the mimicked plan lowering the dose in the right parotid (middle) and their difference in dose distribution (right) for ANON1089. The right parotid is indicated by white arrows.

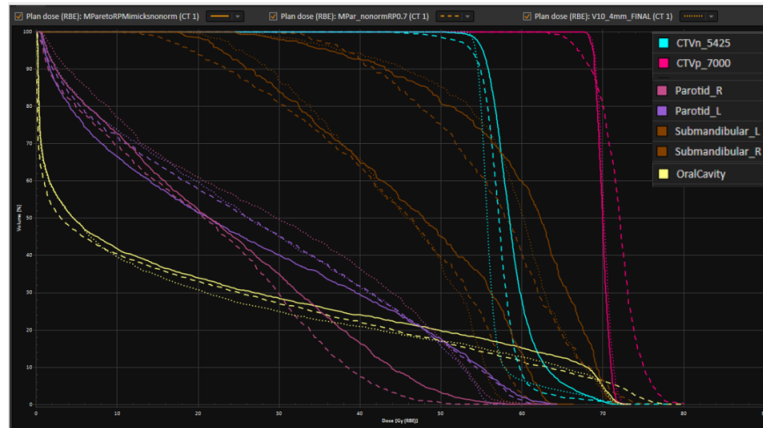


Figure 5.16: Comparison of the DVH curves for the mimicked plan lowering the dose in the right parotid (solid line), the predicted dose lowering the dose in the right parotid (dashed line) and the reference plan (dotted line)

Plan	Dose in Gray [Gy]						
	Oral cavity Dmean	Right parotid Dmean	Left parotid Dmean	Right SMG Dmean	Left SMG Dmean	CTV 5425 D95	CTV 7000 D95
Reference	18.08	28.9	27.10	58.23	43.9	54.3	68.98
Prediction	18.47	20.31	27.09	56.33	43.37	53.80	66.91
Mimicked	19.93	22.54	25.23	58.99	45.63	55.22	68.8

Table 5.1: Dose statistics for the reference, the prediction dose distribution lowering the dose in the right parotid and the mimicked plan for patient ANON1089.

5.2.2 Mimicked plan for a lower dose in the right submandibular gland for patient ANON1089

Figure 5.17 shows the dose distribution comparison between the mimicked plan (middle) and the predicted dose distribution for a lower dose in the right submandibular gland (left) together with their difference in dose (right). The pink color accounts for at least 5 Gy more in the mimicked plan and the light blue for at least 5 Gy less. We observe many spots where the dose is increased by 5 Gy or more in the mimicked plan and particularly around the submandibular glands (the right submandibular gland is indicated by white arrows in Figure 5.17).

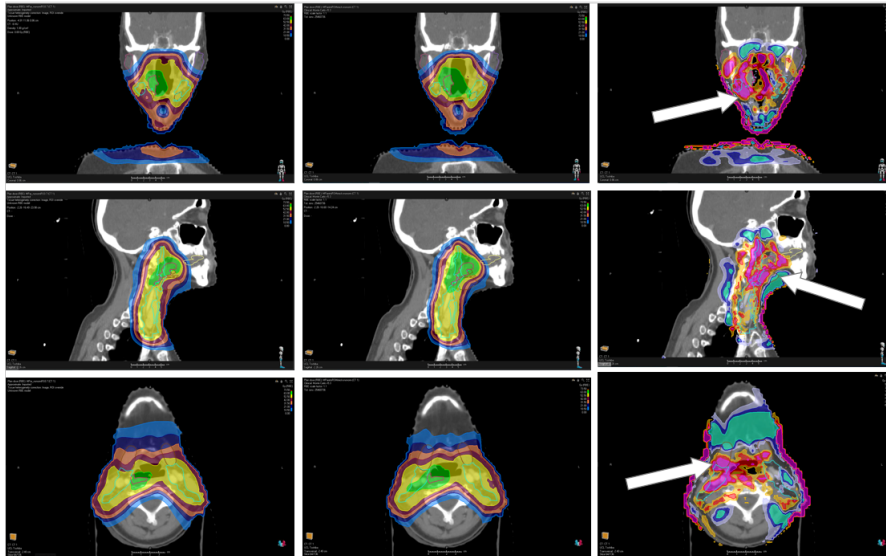


Figure 5.17: Comparison of the predicted dose distribution (left) with the mimicked plan lowering the dose in the right submandibular gland (middle) and their difference in dose distribution (right) for ANON1089. The right submandibular gland is indicated by white arrows.

Figure 5.18 presents the dose distribution comparison between the mimicked plan (middle)

and the reference ground truth dose distribution (left) together with their difference in dose (right). The pink color accounts for at least 5 Gy more in the mimicked plan and the light blue for at least 5 Gy less. We observe many spots where the dose is decreased by 5 Gy or more in the mimicked plan and particularly around the right submandibular gland (indicated by white arrows in Figure 5.18). However, due to the decrease of dose in the right submandibular gland, there are some increases of 5 Gy or more elsewhere.

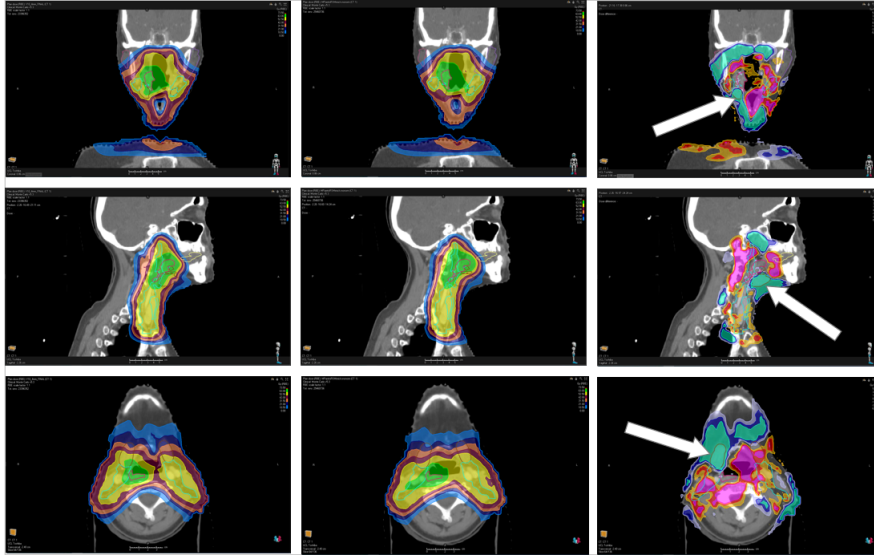


Figure 5.18: Comparison of the ground truth dose distribution (left) with the mimicked plan lowering the dose in the right submandibular gland (middle) and their difference in dose distribution (right) for ANON1089. The right submandibular gland is indicated by white arrows.

Figure 5.19 presents the DVH for some ROIs with a comparison between the mimicked plan (solid line), the predicted dose distribution (dashed line) and the reference plan (dotted line). Some relevant dose volume metrics for the comparison of the three plans are presented in Table 5.2. It is noticeable that the right submandibular has a lower dose for the mimicked plan than for the reference plan with a decrease in mean dose by 9 Gy. However, the mimicked plan has a higher mean dose in the right submandibular gland than the prediction by 4 Gy. The left and right parotid have around the same mean dose in the mimicked and the reference plan. The left submandibular gland and the oral cavity have an increased mean dose of 2 Gy for the mimicked plan compared to the reference plan. The target coverage seems to be as good or even better for the mimicked plan compared to the reference plan.

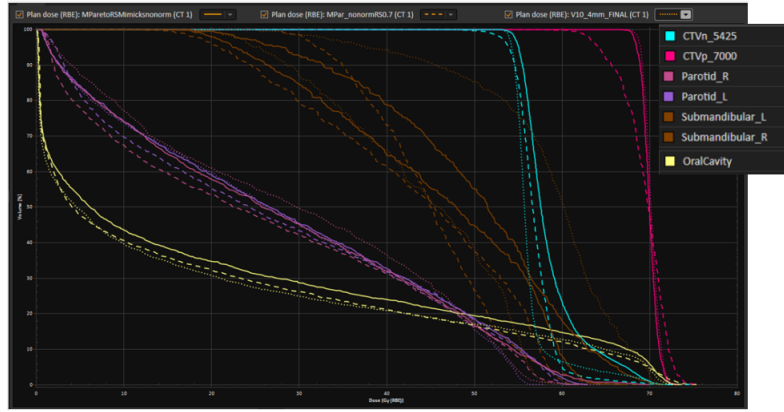


Figure 5.19: Comparison of the DVH curves for the mimicked plan lowering the dose in the right submandibular gland (solid line), the predicted dose lowering the dose in the right submandibular gland (dashed line) and the reference plan (dotted line)

Plan	Dose in Gray [Gy]						CTV 5425 D95	CTV 7000 D95
	Oral cavity Dmean	Right parotid Dmean	Left parotid Dmean	Right SMG Dmean	Left SMG Dmean			
Reference	18.08	28.9	27.10	58.23	43.9		54.3	68.98
Prediction	18.37	25.63	26.85	45	42.99		53.89	64.12
Mimicked	20.29	28.9	27.92	49.56	45.41		55.04	68.51

Table 5.2: Dose statistics for the reference, the prediction dose distribution lowering the dose in the right submandibular gland and the mimicked plan for patient ANON1089.

5.3 Normal Tissue Complication Probability

For each patient, the NTCP percentages for xerostomia of grade 2 and grade 3 are computed for the reference plan, the mimicked plan lowering the dose in the right parotid and the mimicked plan lowering the dose in the right submandibular gland. The results are presented in Table 5.3. Figure 5.20 presents the results in a boxplot to better visualize them.

Between the NTCP of the three plans, small variations of 1 to 3 % can be observed. The mimicked plans lowering the dose in the right parotid have slightly better results than the mimicked plans lowering the dose in the right submandibular gland. This is in fact due to the high dose that the submandibular gland receives even when lowering it. Looking at the boxplots, there is no improvement but no worsening either in the Normal Tissue Complication Probability (NTCP).

Patient	Grade 2 Reference plan	Grade 2 Mimicked RP	Grade 2 Mimicked RSMG	Grade 3 Reference plan	Grade 3 Mimicked RP	Grade 3 Mimicked RSMG
ANON104	38 %	39 %	39 %	10 %	11 %	11 %
ANON1089	42 %	40 %	40 %	11 %	11 %	11 %
ANON1169	53 %	50 %	53 %	16 %	15 %	16 %
ANON1423	37 %	36 %	36 %	10 %	9 %	9 %
ANON1595	46 %	46 %	46 %	13 %	13 %	13 %
ANON1677	39 %	40 %	39 %	10 %	11 %	11 %
ANON2055	36 %	38 %	38 %	10 %	10 %	10 %
ANON2146	37 %	38 %	38 %	10 %	10 %	10 %
ANON2175	42 %	42 %	42 %	12 %	12 %	12 %
ANON2200	40 %	40 %	42 %	11 %	11 %	12 %

Table 5.3: NTCP Xerostomia percentage of grade 2 and 3 for all test patients for the reference plan, the mimicked plan with a lower dose in the right parotid (RP) and the mimicked plan for a lower dose in the right submandibular gland (RSMG)

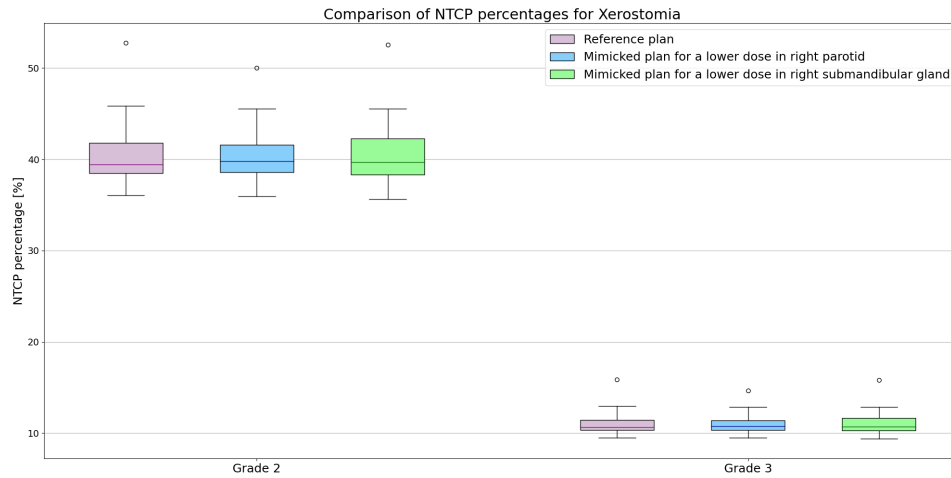


Figure 5.20: Boxplots of the NTCP percentages of Xerostomia of grade 2 and 3 for the reference plan, the mimicked plan with a lower dose in the right parotid and the mimicked plan with a lower dose in the right submandibular gland.

6 | Discussion

In the previous chapter, we analyzed the dose predictions coming from four different models on the test patients. The conventional dose prediction model was first presented to benchmark our results. The second model (the scaled dose prediction model) allowed reducing the dose in an organ with a scaling of 0.7 in the right parotid or the right submandibular gland. The third model is a model that allows changing the dose in the right parotid by a scaling set in the inputs of the model. The last model is a model allowing to change the dose in one organ among the four salivary glands by a scaling between 0.5 and 1.5. Then, a dose mimicking step was performed on the predicted dose distributions from the scaled dose prediction model and the NTCP was computed on the mimicked plans.

This chapter aims at discussing the different results presented in the previous chapter. The limitations of our methods are discussed and compared with existing works on Pareto dose prediction models.

First, the scaled dose prediction models work well to lower the dose in the right parotid or the right submandibular gland. However, these two models were trained with one scaled dose distribution with one scaling of 0.7 only and did not learn how to differentiate a scaling from the reference ground truth. When adding the possibility of choosing the scaling in one organ, more errors occur. The dose is indeed scaled but, not as much as we would desire. The greater the difference of the desired scaling (in an increase or decrease of dose), the greater will be the error. The submandibular glands seem to be more difficult to scale, probably due to their closeness with the CTVs and their very small size compared to the CTV (4% on average). The model understands that the CTVs have a certain prescription level and it is optimized in priority.

The errors may also come from other causes. Some patients do not have submandibular glands. But, more importantly, the patching system is such that a patch of size 96x96x64 of the image is taken while the full image containing the body is usually of size 112x192x96. We make sure the CTV is completely included in the patch but the parotids and submandibular glands might be not included in it. Training the models for more epochs would thus probably allow for a reduction in errors when asking for a scaling far from 1. Another possibility would be to remove patches but it would considerably increase

the training time. Another reason for errors we could think of is that the Pareto model presented by Dan Nguyen et al. in [11] is trained using 1200 plans per patient while we use at most 80 plans per patient for the model scaling a dose in one of the four glands. However, increasing the batch size implies that we increase the training time of the model while it is already 2 days for a batch size of 5 and 100 epochs.

Even if the errors could be reduced, the models presented in this thesis have some limitations. First, they do not take into account variable beam configurations. This could be solved in the same way as in [6, 40] by training a similar model with the desired beam configuration as an input. Secondly, our models only take two different dose prescription levels of 54.25 Gy and 70 Gy as input. A possible improvement would be to train this model for multiple prescription levels.

Additionally, a very simple rather unrealistic approach has been taken for the generation of the Pareto dose distributions needed as ground truths for the training of the models. We have decided to simply scale the dose in a single organ by a given scaling between 0.5 and 1.5. To reach a clinically deliverable plan, dose mimicking was applied to the predicted dose distribution of the scaled dose model. The mimicked plans seem to lower the dose in the desired organ but not as much as desired. A reason for this is that the considered salivary glands are very close to the CTV. The mimicking, as we have defined it, attempts to achieve good coverage of the target volume, mitigating the dose reduction that the predicted dose had achieved. However, mimicking can be defined in other ways by, for example, setting weights for the target isodoses equal to the organs isodoses. But then, the mimicked plan is more likely to be undeliverable and/or not clinically acceptable to a physician. This leads to an additional limitation. Even if the deep learning models take a couple of seconds to predict a dose distribution, the dose mimicking can take up to 20 minutes to generate a robust mimicked plan for proton therapy. For a model predicting many different scalings for different plans, each predicted dose distribution has to be mimicked and the clinical deliverability of each plan has to be checked. This comes down to navigating through all the Pareto mimicked plans a bit in the same way as for MCO. In the end, the time needed to choose the plan accounting for the best trade-offs might be equivalent for the two automatic planning methods.

Furthermore, the NTCP has been computed on the mimicked plans lowering the dose in the right parotid or the right submandibular glands. Their NTCP has been compared to the one for the reference plan for each test patient. Lowering the dose in the right parotid or the right submandibular gland does not seem to have a significant impact on the probability of having xerostomia. This can be explained by the fact that xerostomia is linked to the four salivary glands. And, we have seen in the dose mimicking results that, lowering the dose in one of these four glands, increases the dose level in the three

others. A way to reduce the probability of having xerostomia would thus be to lower the dose in the four glands at the same time. Using a scaled dose distribution with the same scaling of the dose in all four glands seems far from being realistic. Another method needs to be chosen to find the ground truth for the models. A possibility would be to generate Pareto plans through solving MCO problems as did Dan Nguyen et al. [11].

In addition, all existing deep learning Pareto dose prediction models are focused on prostate cancer. This type of cancer has a reduced number of OARs. In this thesis, we focus on head and neck cancer with 13 OARs which are also smaller in volume compared to the CTV volume. The planning of a treatment and the choice of trade-offs are much more difficult than for prostate cancer. However, the Pareto dose prediction is much more interesting as well due to the closeness with the brainstem and other critical OARs. This motivates the use of this simplistic approach of scaling the dose in one organ as a ground truth first instead of directly using Pareto plans generation through MCO problem-solving.

Moreover, the model proposed by Tianfang Zhang et al. [12] was trained to predict a certain level of scaled dose in a certain organ with one model per organ. But, their proposed models did not allow to set the desired scaling as an input. Each of their proposed models thus accounted for one trade-off only. We have proposed a model that allows setting manually which organ we desire to spare and by which scaling. However, our model is still not as specific as the model proposed by Dan Nguyen et al. [11] but it could be improved in that way. The reference ground truth should be modified to a more realistic dose distribution like the ones that could be generated by MCO. As proposed in [12, 39], the DVH loss could be used to help the model to further differentiate the different trade-offs.

Finally, to our knowledge, no deep learning Pareto model has been proposed for proton therapy. They all focus on radiotherapy (IMRT or VMAT). The challenge for proton therapy is that instead of using the PTV to account for uncertainties, we need to ensure that the dose distribution is robust to account for proton range uncertainties and systematic setup errors. In this thesis, the mimicked plan is optimized robustly with some selected robust objective functions. However, we have not assessed the robustness of the mimicked dose.

7 | Conclusion

Because of the potentially numerous side effects associated with radiotherapy treatments, it is important to let the doctor choose the best trade-off for a given patient in the automation of the treatment planning process. Recently, some Pareto dose prediction models have been presented to allow the doctor to set the desired trade-off as an input of the model. This thesis aimed to propose a first proton therapy Pareto dose prediction model with a focus on head and neck cancer.

In this thesis, the notion of Pareto plans and dose trade-offs have been simplified a bit. We have considered scalings of the dose distribution in a single organ as the dose distribution allowing to spare this organ. All models presented in this thesis have been trained using this scaled dose distribution as a ground truth corresponding to a certain trade-off. Three different types of Pareto models have been proposed to adapt the conventional dose prediction model. The first model aimed at predicting a dose distribution with a scaled dose in a single organ with a fixed scaling. The second model aimed at predicting a dose distribution with a scaled dose in a single organ with a scaling chosen between 0.5 and 1.5. The third model aimed at predicting a dose distribution with a scaled dose in one of the four salivary glands with a scaling between 0.5 and 1.5.

The first model is working the best with an average MSE of 1.245 with better results for the model lowering the dose in the right parotid than the submandibular gland. For the two other models, the MSE is 2 to 3 times higher. The more polyvalent the model is, the more the errors increase. In particular, the models have a larger error when asking for a scaling in the submandibular glands. This is partially due to the proximity of these glands to the CTV.

Then, a dose mimicking step has been performed on the dose distributions predicted for the test patients. This step allows to evaluate the deliverability of a lower dose in one of the glands and to have a robust dose distribution. To have a clinically deliverable plan with a lower dose in an organ, the other organs have to see their dose increased. Lowering the dose in only one of the salivary glands thus does not have a significant impact on the NTCP for xerostomia.

The models presented in this thesis are thus not working perfectly and still need to be improved. Some areas for improvement are:

- We should consider a different more realistic Pareto set where the sparing of an organ implies the increase in dose somewhere else. The best would be to take a Pareto set generated through MCO.
- The patches should be removed such that the model can see all organs. If the patches are removed, the DVH loss could also be implemented along the MSE to help the model to further differentiate the different trade-offs.
- The dose mimicking step could be improved by changing the weights of the different objectives of the isodoses. For example, the objectives on the target isodoses might have the same weights as the objectives for the organs isodoses. The resulting mimicked plans should then be analyzed to evaluate their clinical deliverability.

In their improved version, we imagine that the models and methods presented in this thesis could considerably save time in the treatment planning process of a head and neck cancer with proton therapy. In comparison with a conventional dose prediction model, it would also allow the medical physician to choose the best trade-off easily.

Bibliography

- [1] World Health Organization. *Global Health Estimates : Top 10 causes of death in Belgium for both sexes aged all ages (2019)*.
URL: <https://www.who.int/data/gho/data/themes/mortality-and-global-health-estimates/ghe-leading-causes-of-death> (visited on 25/05/2022).
- [2] MedLine Plus. Todd Gersten (MD) and David Zieve (MD). *Cancer treatments*. Oct. 2021.
URL: <https://medlineplus.gov/ency/patientinstructions/000901.htm> (visited on 07/03/2022).
- [3] *Proton Therapy : Advanced Radiation Therapy Treatment*.
URL: <https://provisionhealthcare.com/about-proton-therapy/advantages-of-proton/> (visited on 22/02/2022).
- [4] MJ LaRiviere et al. ‘Proton Therapy’.
In: *Hematol Oncol Clin North Am.* 33.6 (Oct. 2019), pp. 989–1009.
DOI: 10.1016/j.hoc.2019.08.006..
- [5] Ana Maria Barragan Montero.
‘Robust, accurate and patient-specific treatment planning for proton therapy’.
Prom : Lee, John Aldo; Sterpin, Edmond. PhD thesis. UCLouvain, 2017.
- [6] Ana Maria Barragan Montero et al.
‘Three-dimensional dose prediction for lung IMRT patients with deep neural networks: robust learning from heterogeneous beam configurations.’
In: *Med Phys.* 46.8 (Aug. 2019), pp. 3679–3691. DOI: 10.1002/mp.13597..
- [7] Dan Nguyen et al. ‘3D radiotherapy dose prediction on head and neck cancer patients with a hierarchically densely connected U-net deep learning architecture.’
In: *Phys Med Biol.* 64.6 (Mar. 2019), p. 065020.
DOI: 10.1088/1361-6560/ab039b..
- [8] Anni Ravnsbæk Jensen, Hanne Marie Nellemann and Jens Overgaard.
‘Tumor progression in waiting time for radiotherapy in head and neck cancer’.
In: *Radiotherapy and Oncology* 84.1 (2007), pp. 5–10. ISSN: 0167-8140.
DOI: <https://doi.org/10.1016/j.radonc.2007.04.001>. URL: <https://www.sciencedirect.com/science/article/pii/S0167814007001533>.

- [9] R. M. Wyatt, A. H. Beddoe and R. G. Dale.
‘The effects of delays in radiotherapy treatment on tumour control’.
In: *Physics in Medicine and Biology* 48.2 (2003), pp. 139–155.
DOI: 10.1088/0031-9155/48/2/301..
- [10] Elsevier. *Pareto Optimality - an overview / ScienceDirect Topics*. URL:
<https://www.sciencedirect.com/topics/engineering/pareto-optimality>
(visited on 01/06/2022).
- [11] Dan Nguyen et al. ‘Generating Pareto optimal dose distributions for radiation therapy treatment planning’. In: (2019). DOI: 10.48550/ARXIV.1906.04778.
URL: <https://arxiv.org/abs/1906.04778>.
- [12] Tianfang Zhang, Rasmus Bokrantz and Jimmy Olsson. ‘Probabilistic Pareto plan generation for semiautomated multicriteria radiation therapy treatment planning’. In: *Physics in Medicine and Biology* 67.4 (Feb. 2022), p. 045001. ISSN: 1361-6560.
DOI: 10.1088/1361-6560/ac4da5.
URL: <http://dx.doi.org/10.1088/1361-6560/ac4da5>.
- [13] Jianhui Ma et al. ‘A feasibility study on deep learning-based individualized 3D dose distribution prediction’. In: *Medical Physics* 48.8 (July 2021), pp. 4438–4447.
DOI: 10.1002/mp.15025. URL: <https://doi.org/10.1002%5C%2Fmp.15025>.
- [14] Andy Trotti. ‘Toxicity in head and neck cancer: a review of trends and issues’. In: *International Journal of Radiation Oncology*Biophysics* 47.1 (2000), pp. 1–12. ISSN: 0360-3016.
DOI: [https://doi.org/10.1016/S0360-3016\(99\)00558-1](https://doi.org/10.1016/S0360-3016(99)00558-1). URL:
<https://www.sciencedirect.com/science/article/pii/S0360301699005581>.
- [15] Piet Dirix, Sandra Nuyts and Walter Van den Bogaert.
‘Radiation-induced xerostomia in patients with head and neck cancer’.
In: *Cancer* 107.11 (2006), pp. 2525–2534.
DOI: <https://doi.org/10.1002/cncr.22302>. eprint: <https://acsjournals.onlinelibrary.wiley.com/doi/pdf/10.1002/cncr.22302>.
URL: <https://acsjournals.onlinelibrary.wiley.com/doi/abs/10.1002/cncr.22302>.
- [16] Rasmus Bokrantz. ‘Multicriteria optimization for managing tradeoffs in radiation therapy treatment planning’. Prom : Forsgren, Anders.
PhD thesis. KTH Royal Institute of Technology, 2013.
- [17] National Institute of Biomedical Imaging and Bioengineering.
Computed Tomography (CT). URL: <https://www.nibib.nih.gov/science-education/science-topics/computed-tomography-ct> (visited on 19/04/2022).

- [18] Minh Tam Truong and Nataliya Kovalchuk. ‘Radiotherapy Planning’. In: *PET Clinics* 10.2 (2015). PET/CT and Patient Outcomes, Part I, pp. 279–296. ISSN: 1556-8598. DOI: <https://doi.org/10.1016/j.cpet.2014.12.010>. URL: <https://www.sciencedirect.com/science/article/pii/S1556859814001424>.
- [19] Charlotte L. Brouwer et al. ‘CT-based delineation of organs at risk in the head and neck region: DAHANCA, EORTC, GORTEC, HKNPCSG, NCIC CTG, NCRI, NRG Oncology and TROG consensus guidelines’. In: *Radiotherapy and Oncology* 117.1 (2015), pp. 83–90. ISSN: 0167-8140. DOI: <https://doi.org/10.1016/j.radonc.2015.07.041>. URL: <https://www.sciencedirect.com/science/article/pii/S0167814015004016>.
- [20] Hakan Nystrom, Maria Jensen and Petra Nyström. ‘Treatment planning for Proton therapy: What is needed in the next 10 years?’. In: *The British Journal of Radiology* 93 (July 2019), p. 20190304. DOI: 10.1259/bjr.20190304.
- [21] Neil G. Burnet. ‘Defining the tumour and target volumes for radiotherapy’. In: *Cancer Imaging* 4.2 (2004), pp. 153–161. DOI: 10.1102/1470-7330.2004.0054. URL: <https://doi.org/10.1102/1470-7330.2004.0054>.
- [22] Hisham Kamal-Sayed et al. ‘Adaptive method for multicriteria optimization of intensity-modulated proton therapy’. In: *Medical Physics* 45.12 (2018), pp. 5643–5652. DOI: <https://doi.org/10.1002/mp.13239>. eprint: <https://aapm.onlinelibrary.wiley.com/doi/pdf/10.1002/mp.13239>. URL: <https://aapm.onlinelibrary.wiley.com/doi/abs/10.1002/mp.13239>.
- [23] Dan Nguyen et al. *Generating Pareto optimal dose distributions for radiation therapy treatment planning*. 2019. DOI: 10.48550/ARXIV.1906.04778. URL: <https://arxiv.org/abs/1906.04778>.
- [24] Chen Wei et al. ‘A fast optimization algorithm for multicriteria intensity modulated proton therapy planning’. In: *Medical physics* 37 9.9 (Sept. 2010), pp. 4938–45. ISSN: 0094-2405. DOI: 10.1118/1.3481566. URL: <https://www.osti.gov/biblio/22100588>.
- [25] S E McGowan, N G Burnet and A J Lomax. ‘Treatment planning optimisation in proton therapy’. In: *The British Journal of Radiology* 86.1021 (2013). PMID: 23255545, pp. 20120288–20120288. DOI: 10.1259/bjr.20120288. eprint: <https://doi.org/10.1259/bjr.20120288>. URL: <https://doi.org/10.1259/bjr.20120288>.

- [26] Joseph K Kim et al. ‘Proton Therapy for Head and Neck Cancer’.
In: *Current treatment options in oncology* 19.6 (May 2018), p. 28. ISSN: 1527-2729.
DOI: 10.1007/s11864-018-0546-9.
URL: <https://doi.org/10.1007/s11864-018-0546-9>.
- [27] Wei Chen et al. ‘Including Robustness in Multi-criteria Optimization for Intensity Modulated Proton Therapy’.
In: *Physics in medicine and biology* 57 (Feb. 2012), pp. 591–608.
DOI: 10.1088/0031-9155/57/3/591.
- [28] Marjan Sharabiani et al. ‘Generalizability assessment of head and neck cancer NTCP models based on the TRIPOD criteria’.
In: *Radiotherapy and Oncology* 146 (2020), pp. 143–150. ISSN: 0167-8140.
DOI: <https://doi.org/10.1016/j.radonc.2020.02.013>. URL: <https://www.sciencedirect.com/science/article/pii/S0167814020300761>.
- [29] Ana Maria Barragan Montero et al. ‘Artificial intelligence and machine learning for medical imaging: A technology review’.
In: *Physica Medica* 83 (Mar. 2021), pp. 242–256.
DOI: 10.1016/j.ejmp.2021.04.016.
- [30] June-Goo Lee et al. ‘Deep Learning in Medical Imaging: General Overview’.
In: *Korean Journal of Radiology* 18 (July 2017), p. 570.
DOI: 10.3348/kjr.2017.18.4.570.
- [31] Geoff Currie et al. ‘Machine learning and deep learning in medical imaging: Intelligent imaging’. English. In: *Journal of Medical Imaging and Radiation Sciences* 50.4 (Dec. 2019). this is original research by definition. it creates learning and new knowledge in the context of the readership, pp. 477–487. ISSN: 1876-7982.
DOI: 10.1016/j.jmir.2019.09.005.
- [32] Maciej A Mazurowski. ‘Deep learning in radiology: An overview of the concepts and a survey of the state of the art with focus on MRI.’
In: *Journal of magnetic resonance imaging*. 49.4 (Dec. 2018). ISSN: 1053-1807.
- [33] Rikiya Yamashita et al.
‘Convolutional neural networks: an overview and application in radiology’.
In: *Insights into Imaging* 9 (June 2018). DOI: 10.1007/s13244-018-0639-9.
- [34] Yann LeCun, Y. Bengio and Geoffrey Hinton. ‘Deep Learning’.
In: *Nature* 521 (May 2015), pp. 436–44. DOI: 10.1038/nature14539.
- [35] Olaf Ronneberger, Philipp Fischer and Thomas Brox.
‘U-Net: Convolutional Networks for Biomedical Image Segmentation’.
In: *CoRR* abs/1505.04597 (2015). arXiv: 1505.04597.
URL: <http://arxiv.org/abs/1505.04597>.

- [36] Antonin Chambolle and Thomas Pock. ‘A First-Order Primal-Dual Algorithm for Convex Problems with Applications to Imaging’.
In: *Journal of Mathematical Imaging and Vision* 40 (2010), pp. 120–145.
- [37] RaySearch Laboratories. *Advanced proton planning in RayStation with robust optimization and Multi-Criteria Optimization*. 2015.
URL: <https://www.youtube.com/watch?v=39u2w-Nni1w> (visited on 13/05/2022).
- [38] Wentao Wang et al. ‘Fluence Map Prediction Using Deep Learning Models – Direct Plan Generation for Pancreas Stereotactic Body Radiation Therapy’.
In: *Frontiers in Artificial Intelligence* 3 (2020). DOI: 10.3389/frai.2020.00068.
- [39] Dan Nguyen et al.
‘Incorporating human and learned domain knowledge into training deep neural networks: A differentiable dose-volume histogram and adversarial inspired framework for generating Pareto optimal dose distributions in radiation therapy’.
In: *Med Phys* 47.3 (Mar. 2020), pp. 837–849. DOI: 10.1002/mp.13955.
- [40] Gyanendra Bohara et al. ‘Using deep learning to predict beam-tunable Pareto optimal dose distribution for intensity-modulated radiation therapy’.
In: *Medical Physics* 47.9 (Aug. 2020), pp. 3898–3912. DOI: 10.1002/mp.14374.
URL: <https://doi.org/10.1002/mp.14374>.
- [41] P. James Jensen et al.
‘A Novel Machine Learning Model for Dose Prediction in Prostate Volumetric Modulated Arc Therapy Using Output Initialization and Optimization Priorities’.
In: *Frontiers in Artificial Intelligence* 4 (2021). ISSN: 2624-8212.
DOI: 10.3389/frai.2021.624038.
URL: <https://www.frontiersin.org/article/10.3389/frai.2021.624038>.
- [42] Keras. *Adam*.
URL: <https://keras.io/api/optimizers/adam/> (visited on 23/05/2022).
- [43] *Landelijk Indicatie Protocol Protonentherapie (versie 2.2) (LIPPv2.2) Hoofd-halstumoren*.
Nederlandse Vereniging voor Radiotherapie en Oncologie, 2019.
URL: <https://nvro.nl/publicaties/rapporten> (visited on 25/05/2022).

UNIVERSITÉ CATHOLIQUE DE LOUVAIN
École polytechnique de Louvain

Rue Archimède, 1 bte L6.11.01, 1348 Louvain-la-Neuve, Belgique | www.uclouvain.be/epl