

L'Institut Agro Rennes-Angers
☐ Site d'Angers ☒ Site de Rennes

Année universitaire : 2024-2025
Spécialité :
Ingénieur agronome
Spécialisation (et option éventuelle) :
Sciences halieutiques et aquacoles

Mémoire de fin d'études

- ☒ d'ingénieur de l'Institut Agro Rennes-Angers (Institut national d'enseignement supérieur pour l'agriculture, l'alimentation et l'environnement)
☐ de master de l'Institut Agro Rennes-Angers (Institut national d'enseignement supérieur pour l'agriculture, l'alimentation et l'environnement)
☐ de l'Institut Agro Montpellier (étudiant arrivé en M2)
☐ d'un autre établissement (étudiant arrivé en M2)

Identification des pathologies sur photos de migrateurs amphihalins par analyse d'images

Par : Titouan LECOMTE



Soutenu à Rennes, le 11/09/2025

Devant le jury composé de :

Président : Etienne Rivot
Maître de stage : François Martignac, Jérôme Guitton,
Marie-Pierre Etienne

Autres membres du jury : Pablo Brosset (Enseignant-chercheur L'Institut Agro Rennes-Angers) et Laetitia Chapel (Enseignante-chercheuse L'Institut Agro Rennes-Angers)

Les analyses et les conclusions de ce travail d'étudiant n'engagent que la responsabilité de son auteur et non celle de l'Institut Agro Rennes-Angers

Ce document est soumis aux conditions d'utilisation «Patrimoine-Pas d'Utilisation Commerciale-Pas de Modification 4.0 France» disponible en ligne <http://creativecommons.org/licenses/by-nc-nd/4.0/deed.fr>



Remerciements

Je tiens à remercier chaleureusement François Martignac, Jérôme Guitton et Marie-Pierre Etienne qui m'ont encadré et accompagné durant ce stage. Leur approche du sujet pour me permettre de m'approprier un sujet très nouveau m'a permis de prendre en main les outils d'intelligences artificielles d'analyse d'images pour la suite du projet. La bienveillance et la complémentarité de cet encadrement m'ont permis de mettre en place chaque étape du pipeline d'obtention des codes patho et les stratégies d'évaluations des performances. Je tenais à mentionner leur aide lors des opportunités que j'ai eu de présenter mon stage à l'AG de DECOD, à la journée Amédée ou encore lors des présentations aux opérateurs et aux journées techniques de l'ORE, convié par Frédéric Marchand.

Je remercie également Aziliz Floquet, Quentin Josset, Nicolas Jeannot, Fabien Quendo et Julien Tremblay qui se reconnaîtront dans ce rapport au travers de nombreuses mentions en tant qu'« opérateurs » ou encore « experts ». Ils se sont mobilisés pour nous fournir des données de codes pathologie et d'images standardisées et n'ont pas hésité à donner de leur temps pour participer à la labellisation des données d'apprentissage. Leurs ressources bibliographiques et les différents échanges sur l'opérationnalité du projet m'ont permis de construire le stage selon leurs attentes.

Un remerciement également aux doctorants Baptiste, Emilien et Emilio qui m'ont accueilli dans leur bureau durant tout le stage, sans qui, mes pauses n'auraient pas été aussi ressourçantes.

Table des matières

Introduction.....	1
1. Contexte	1
2. Etat de l'art	3
Matériel et méthodes	8
1. Données	8
2. Modèles.....	11
3. Construction du jeu de données et apprentissage	11
a. Labellisation	11
5. Estimation de la localisation	13
6. Evaluation des performances du pipeline	15
7. Application à un cas concret	18
8. Déroulé du pipeline	18
Résultats	19
2. Evaluation de l'apprentissage du modèle de détection des pathologies	21
a. Sélection du meilleur modèle	21
3. Evaluation des prédictions de la gravité	22
4. Evaluation des prédictions de la localisation.....	24
5. Evaluation générale du modèle et le pipeline	25
6. Précision, sensibilité et spécificité du modèle choisi.....	27
7. Exemple d'application	28
Discussion	30
1. Phase de labellisation	30
2. Partitionnement du poisson.....	30
3. Evaluation du modèle de détection	31
4. Evaluation de la prédiction du code patho.....	31
5. Choix du seuil de confiance.....	32
6. Précision, sensibilité et spécificité du modèle de prédiction de présence/absence	33
7. Application à un cas concret	33
Conclusion	34
Perspectives.....	35
Bibliographie	36
Annexes	41

Liste des figures

Figure 1 : Interface de saisie des pathologies sur le poste de biométrie	4
Figure 2 : a) Schéma de l'organisation du Deep-Learning. b) Schéma explicatif (CNN, © MapR, C.D, Futura)	5
Figure 3 : Performances des modèles Yolo	8
Figure 4 : Exemple d'image acquise par les opérateurs	8
Figure 5 : Saisonnalité de l'occurrence des pathologies et nombre de poissons rencontrés	11
Figure 6: Extrait du schéma du pipeline (b. en annexe V) d'extraction des codes pathos d'un individu	19
Figure 7 : Distribution des erreurs de mesures de la longueur fourche (% de la longueur de référence)	20
Figure 8 : Relations entre la hauteur totale du poisson au niveau de la tête et la hauteur de la ligne latérale (a.) et entre la hauteur totale du poisson au niveau du pédoncule caudal et la hauteur de la ligne latérale (b.).....	20
Figure 9 : Partitionnement morphologique automatisé du poisson	21
Figure 10 : Matrice de confusion des détections du modèle sélectionné.....	22
Figure 11 : Performances de prédiction de l'indice de gravité du modèle pour les pertes d'écailles ...	23
Figure 12 : Performances de prédiction de l'indice de gravité du modèle pour les rougeurs	24
Figure 13 : Performances de prédiction de l'indice de gravité du modèle pour les poux.....	24
Figure 14 : Performances de prédiction du code de localisation du modèle pour les trois pathologies	25
Figure 15 : Performances de prédiction de localisation et de gravité pour les pertes d'écailles.....	26
Figure 16 : Performances de prédiction de localisation et de gravité pour les rougeurs.....	26
Figure 17 : Performances de prédiction de localisation et de gravité pour les poux	26
Figure 18 : Exemple de détection des poux sur le corps d'un poisson	28
Figure 19 : Comparaison de l'estimation de la proportion d'individus infectés par les poux aux valeurs observées	29
Figure 20 : Comparaison des histogrammes du nombre d'individus atteints par les poux par semaine (a) observés, (b) prédits	29

Liste des tableaux

Tableau 1 : Détail des indices de gravité des codes patho (Elie, Girard 2014)	3
Tableau 2 : Nombre d'individus disponibles pour notre étude par année et par espèce	9
Tableau 3 : Extrait des codes relatifs aux pathologies.....	10
Tableau 4 : Extrait des codes relatifs aux localisations des pathologies	10
Tableau 5 : Nombre d'occurrence des pathologies utilisées lors de l'apprentissage de Yolo	13
Tableau 6 : Performances du modèle par pathologie aux seuils de confiance sélectionnés.....	27

Liste des annexes

Annexe I : Squelette du masque du poisson	41
Annexe II : Démarche d'obtention du point fourche	42
Annexe III: Courbe de régression du tracé du squelette	43
Annexe IV : Arbre de décision des codes de localisation.....	44
Annexe V : Schéma global de la démarche du stage : entraînement du modèle de détection des pathologies, pipeline et évaluation de la méthode	45
Annexe VI : Valeur de perte des prédictions des boîtes au cours des apprentissages dans le sweep ..	46
Annexe VII : Valeur de mAP50-95 au cours des apprentissages dans le sweep	47

Introduction

1. Contexte

Ces dernières décennies, on observe, sur le territoire français et européen, un déclin des espèces aquatiques migratrices diadromes. Ce constat est souvent illustré par la forte diminution voire l'arrêt de l'activité de pêche du saumon atlantique (*Salmo salar*), professionnelle et récréative (ICES 2024a; DIRM - Manche Est-Mer du Nord 2025). De par la nature de leur cycle de vie, ces animaux migrateurs rencontrent, dans les eaux continentales et marines, de nombreuses perturbations et pressions d'origines climatique et anthropique (Merg et al. 2020).

Des stations de contrôle des poissons migrateurs sont implantées sur les cours d'eau français et sont le lieu de suivi des espèces aquatiques migratrices, et notamment amphihalines. Ces structures sont sous la tutelle de l'OFB (Office Français de la Biodiversité) au travers, notamment, de l'ORE DiaPFC (Observatoire de Recherche sur les Poissons Diadromes dans les Fleuves Côtiers). Présents sur la Nivelle (Pays Basque), le Scorff (Bretagne), l'Oir et la Bresle (Normandie), les opérateurs sur place capturent les poissons des espèces étudiées lors de leur passage au niveau des pièges lors de leur migration, à la remontée en eau douce pour les individus adultes. Ces manipulations ont pour objectif la collecte d'échantillons biologiques, de données morphologiques (taille, poids) et l'évaluation de leur état de santé en notifiant la présence de pathologies externes.

L'enregistrement de ces pathologies est important pour le suivi de l'état de santé des populations de migrateurs en milieu naturel et pour connaître les tendances spatiales et temporelles de ces maladies. En effet, dès 1980, de nombreux travaux se sont basés sur la santé des individus comme indice d'intégrité biologique des milieux, au même titre que les indicateurs de biodiversité taxonomique et fonctionnelle (Karr 1981; Fausch, Karr, Yant 1984).

Dans le cadre de notre étude, on retrouve entre autres les plaies causées par des filets de pêcheurs, morsures de phoques, loutres ou autres prédateurs, les organes abîmés ou manquants, des pétéchies (rougeur cutanée, souvent associée à une anémie infectieuse (Hårstein, Bullock 1976; Bakke, Harris 1998)), ou encore la présence de parasites externes comme des poux de mer, principalement *Lepeophtheirus salmonis* sur les salmonidés. Ces petits crustacés (copépodes) sont à l'origine de chocs osmotiques, infections et transmissions de virus et bactéries aux salmonidés infectés (Nylund, Bjørknes, Wallace 1991). Ces invertébrés sont au cœur des préoccupations actuelles quant à l'élevage intensif de saumon dans les eaux nordiques, à l'origine de foyers d'épidémies de parasitismes (Costello 2009). Ces transmissions de parasites externes entraînent un affaiblissement des populations sauvages, voire un retour prématuré en eau douce. Ainsi, le succès reproducteur des poissons infectés est diminué par rapport à un individu sain (Bolstad et al. 2025).

Bien que théorisées par un manuel reprenant l'historique de ces pratiques à l'échelle nationale, les méthodes et les applications (Elie, Girard 2014), de telles données sur la santé des poissons restent, en pratique, très subjectives, car elles sont dépendantes de l'appréciation d'un observateur, à un instant précis. Par exemple, après une longue série d'observations de pathologies d'une importance élevée, le jugement peut être porté vers une gravité sous-estimée pour les enregistrements suivants.

De même, sur un poisson très affecté par une ou plusieurs pathologies majoritaires, une minoritaire est susceptible d'être omise. A cela s'ajoute l'interprétation de chaque expert des symptômes exprimés par l'individu examiné, sans avoir le temps d'approfondir le diagnostic par des analyses allant au-delà du contrôle visuel, le temps hors d'eau des individus étant minimisé pour améliorer le bien-être animal.

En psychologie, ces phénomènes sont bien documentés, détaillant la manifestation de biais et leurs conséquences sur la mémoire et le jugement. Entre autre, le biais de position sérielle, avec l'effet de primauté et de récence, augmentant l'attention d'un sujet au début et à la fin d'une série de tâches, très étudiée en milieu judiciaire, durant les procès (Ebbinghaus 2013; Enescu 2009). Le biais de contraste décrit, quant à lui, la possible influence de la perception d'un élément par rapport au précédent (Wegener, Petty 1995). Au total, en imagerie médicale, selon une étude sur 484 cas, 53 % des diagnostics comportaient des erreurs d'origine cognitive et 34 % impliquaient une erreur de perception (Taylor et al. 2011). De même, en ressources humaines, des processus d'ancrage, expliquant l'influence des premiers cas sur la perception des suivants, sont mis en évidence chez les auditeurs (Joyce, Biddle 1981; Blondiaux et al. 2018). Enfin, un autre biais pourrait entrer en jeu lors de l'habituation de l'opérateur à émettre une longue série de jugements (Davelaar et al. 2011). C'est pourquoi, dans les domaines de la médecine animale, opérateurs et vétérinaires se concertent pour homogénéiser les diagnostics (Bertram, Klopfleisch 2017; Egevad et al. 2017; Rey et al. 2020).

Ces dernières années, l'explosion du développement et la démocratisation d'outils d'analyse automatique d'images par intelligence artificielle visent notamment à répondre à ces problématiques. Ces méthodes émergentes, appuyées par les structures de réseaux de neurones convolutifs (CNN), permettent en effet d'identifier des objets automatiquement sur une image. Ainsi, ces algorithmes pourraient être en mesure de répondre aux attentes des opérateurs des stations de piégeage quant à l'automatisation de l'acquisition des données pathologiques. La détection et l'identification des dommages et parasites sur des photographies de poissons capturés sur place, ainsi que l'estimation de la surface occupée par la pathologie seraient déterminés par une source unique, consensus des experts. La normalisation de ce diagnostic pourrait apporter une solution d'analyse instantanée de l'état de l'individu, apportant une cartographie des dommages rapide et applicable à tous les individus, tâche que les opérateurs pourront superviser pour s'assurer de la conformité des codes obtenus. Cette volonté de standardisation des méthodes de diagnostic et collecte des données est une problématique à laquelle la communauté scientifique travaille depuis des décennies pour uniformiser la marche à suivre et la nature des données à intégrer, en milieu marin comme dulçaquicole (Belpaire, Goemans 2007; Bucke et al. 1996).

Dans ce contexte, le stage que je réalise à l'Institut Agro, en collaboration avec l'INRAE, au sein de l'UMR DECOD vise à évaluer le potentiel de ces nouveaux outils d'analyse d'images appliqués à ces problématiques écologiques d'évaluation de la santé des stocks ichthyologiques. Pour ce faire, il faudra détecter les pathologies sur chaque image, localiser leur position sur le corps du poisson et quantifier leur gravité en fonction de leur nombre ou de la surface que représente les dommages par rapport à la surface totale du corps de l'individu. De plus, pour connaître l'efficacité de cette méthode à estimer le diagnostic des experts, nous chercherons à évaluer les performances de la démarche sur sa capacité

à prédire une présence ou absence et son erreur d'estimation de la gravité et de la localisation de chaque pathologie.

Cette étude portera exclusivement sur deux espèces migratrices amphihalines rencontrées dans les stations de piégeages et au cœur des préoccupations actuelles de conservation (ICES 2024b) : le saumon atlantique, *Salmo salar* et la truite de mer, *Salmo trutta*. Ces deux espèces du genre *Salmo* sont morphologiquement proches et peuvent être considérées comme un groupe uniforme d'individus semblables pour l'analyse d'images en apprentissage profond (*deep-learning*). Ces espèces migratrices anadromes naissent en eau douce, avant de migrer vers l'océan durant le printemps, quelques années après leur éclosion. Leur retour en eau douce s'effectue au bout d'un ou deux ans généralement, durant le deuxième et troisième trimestre pour se reproduire en hiver (Jensen 1968; Youngson et al. 1983).

2. Etat de l'art

a. Les codes pathologie

Les informations concernant les pathologies externes sont bancarisées en utilisant des codes spécifiques, les codes pathologie, usuellement abrégés « codes patho ». Chaque code est composé de trois informations (Elie, Girard 2007; 2014):

- Deux lettres spécifiques à la nature de la pathologie (cf. Figure 1).
- Deux lettres relatives à la localisation de la pathologie sur le corps du poisson (cf. Figure 1).
- D'un nombre entier compris entre 0 et 4, relatant de la gravité de la pathologie décrite, sur un critère de taux de recouvrement de la maladie ou nombre d'occurrences d'un parasite sur l'ensemble du corps de l'individu (cf. Tableau 1).

Il existe également un code spécifique indiquant qu'aucune pathologie n'a été identifiée par l'opérateur : OOCT0. Ce qui est à distinguer des cas d'absence d'information concernant les codes patho, identifiés par des NA dans la base.

Tableau 1 : Détail des indices de gravité des codes patho (Elie, Girard 2014)

Code gravité	Nombre de lésion / abondance parasitaire, N	Taux de recouvrement corporel, S2 (%)
0	N = 0 ou absence	S2 = 0, absence
1	N < 3, abondance faible	S2 < 5
2	3 < N < 7, abondance moyenne	5 ≤ S ² < 10
3	7 ≤ N < 10, abondance forte	10 ≤ S ² ≤ 20
4	N ≥ 10, abondance très forte	S ² ≥ 20

Les codes de localisation des pathologies sont normalisés et contiennent des unités morphologiques et des agrégations d'unités comme « CT » qui correspond au poisson entier. Tous ces codes sont mis en forme sur l'interface utilisée par les opérateurs sur le terrain pour l'acquisition de données. On retrouve un schéma de poisson fusiforme et un anguilliforme sur la table de biométrie, et l'opérateur doit cliquer sur la partie atteinte par la pathologie pour entrer l'information (cf. Figure 1).

Ce protocole d'acquisition des informations concernant les pathologies externes est aujourd'hui en réflexion par les opérateurs afin d'expérimenter une simplification du niveau de détail attendu (Josset,

comm. pers.). On peut citer la pathologie « rougeur » qui regroupe notamment les pétéchies et hémorragies cutanées. Ce regroupement de certains codes a pour objectif de simplifier les études portant sur les pathologies des poissons anadromes.

Figure 1 : Interface de saisie des pathologies sur le poste de biométrie

b. L'analyse d'image par réseau de neurones convolutif profond

i. Présentation

L'analyse d'images par des méthodes statistiques d'apprentissage profonds a connu ces dernières années, une montée en puissance et une démocratisation au sein de la communauté scientifique, notamment porté par l'apparition des modèles YOLO, You Only Look Once (Redmon et al. 2016). Ces modèles sont devenus populaires dans le monde scientifique car ils permettent de détecter et reconnaître des objets sur une image de manière automatique. La structure de ces modèles est basée sur une architecture de réseaux de neurones convolutifs (Lecun et al. 1998), inspirée du fonctionnement du cerveau et fait référence à des méthodes d'intelligences artificielles (cf. Figure 2.a). Ce réseau est composé de plusieurs couches de neurones virtuels, qui apprennent à repérer des entités visuelles automatiquement grâce à l'extraction de caractéristiques de ces objets comme la netteté, les textures, les contours etc. Chaque neurone est décrit par une combinaison linéaire des sorties des neurones de la couche précédente. Cette efficacité innovante se base sur une phase d'apprentissage, durant laquelle le modèle analyse un grand nombre d'images pour ajuster ses paramètres et extraire au mieux les composantes visuelles des objets à détecter pour minimiser l'erreur de prédiction. Par la suite, le modèle est en mesure de détecter les objets sur de nouvelles images (cf. Figure 2.b), dont les performances sont quantifiées lors de la validation du modèle.

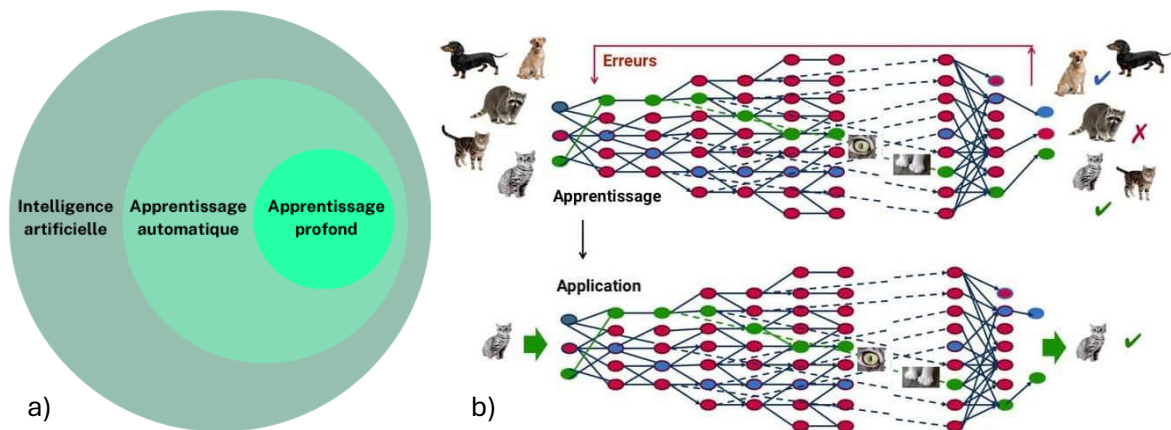


Figure 2 : a) Schéma de l'organisation du Deep-Learning. b) Schéma explicatif (CNN, © MapR, C.D, Futura)

Cette phase d'entraînement est dite supervisée, car elle nécessite un jeu de données « labellisées ». La labellisation consiste à informer la localisation sur l'image et la classe des objets d'intérêts, qui serviront de référence pour ajuster l'apprentissage et valider le modèle. Lors de l'entraînement, le jeu de données, composé des images et leurs labels associés, est subdivisé en trois parties : la plus conséquente, pour l'entraînement du modèle d'environ 70 % des images, puis une part servira à la validation du modèle de 10 % et à ajuster ses paramètres lors d'un prochain cycle d'apprentissage (i.e. époque) et enfin la dernière partie de test, constitué des 20 % de données restantes, évaluera le modèle avec des données indépendantes à l'entraînement.

Le processus d'apprentissage a pour objectif de minimiser une fonction de perte, qui quantifie l'écart entre les estimations et les observations, à l'aide d'un algorithme d'optimisation. Au cours des itérations, ces pertes d'informations devraient être minimisées avant d'atteindre un plateau. Au-delà, le modèle risque un sur-apprentissage (ou *over-fitting*) sur les données, s'opposant à une généralisation de ses capacités de détection. Ces mécanismes sont souvent évités par des stratégies de régularisation (Neyshabur et al. 2017).

ii. Les hyperparamètres et modalités d'apprentissage

Les modalités de cet apprentissage sont décrites, entre autres, par les hyperparamètres de l'entraînement. Qu'ils soient qualitatifs, comme la nomination d'une fonction spécifique ou quantitatifs, les modèles Yolo peuvent tenir compte d'un grand nombre de ces modalités, tels que :

- Epoch : nombre d'itérations (i.e. époques) de la phase d'apprentissage. Trop peu d'époques ne permettent pas au modèle de s'ajuster, mais un trop grand nombre laisserait place à l'*over-fitting* et un temps de calcul important, voire contre-productif.
- Batch_size : nombre d'images traitées avant que le modèle ne mette à jour ses poids. Plus ce lot est grand, plus l'entraînement est rapide et se généralise, mais dans le cas de fortes variabilités au sein d'une même classe (objet diffus comme les tâches), un lot d'images limité est préférable, car le signal moyen du lot sera moins bruité et l'analyse sera plus fine.

- **Image_size** : dimensions de l'image en pixels. Plus la résolution de l'image est élevée, plus le modèle peut extraire de détails dans les patterns des objets, mais une qualité moindre simplifie le modèle et rend l'entraînement plus rapide, nécessitant moins de mémoire.
- **Initial_lr** : *learning rate* initial, conditionnant l'amplitude de la mise à jour des poids du modèle pour l'optimisation de la fonction de perte. Une faible valeur peut entraîner un apprentissage lent et bloqué dans un minimum local de perte, tandis qu'une valeur trop élevée peut faire diverger l'apprentissage, sans avoir rencontré de minimum.
- **Weight_decay** : pénalisation des poids du modèle, pénalisant l'attribution de poids trop grands, conduisant à un over-fitting sur certains patterns des images. Il favorise donc la généralisation de la détection.

Des combinaisons différentes de ces hyperparamètres vont entraîner des performances diverses et peuvent être optimisées par des algorithmes de balayage (ou sweep), tels que Weight and Biases (W&B 2025). Cela consiste à explorer l'espace des hyperparamètres et tester l'effet de chacun sur les métriques de performances d'apprentissage du modèle pour l'optimiser. Cette méthode produit donc de nombreux modèles, dont les modalités d'apprentissage sont différentes, pour en extraire la meilleure combinaison testée. Ce choix des hyperparamètres est un levier central pour améliorer les performances de l'apprentissage du modèle de détection.

iii. L'augmentation de données

Dans de nombreuses études, le jeu de données est limité et ne permet pas une généralisation de l'extraction des caractéristiques de objets. C'est pour combler ce manque de variabilité qu'il existe des méthodes d'augmentation de données. En analyse d'images, cela consiste à modifier sensiblement les images par des rotations, translations, changements d'échelle, rognage, modifications de la luminosité, du contraste ou de la netteté de l'image ou encore l'extraction des contours des objets de l'image (i.e. edge), correspondant aux limites de transitions fortes d'intensité sur une image. A chaque effet est associé une probabilité d'être appliqué et une intensité d'application. Ainsi, à partir d'une image, on peut en créer artificiellement plusieurs autres afin d'augmenter la variabilité. Selon les degrés de modifications et le nombre de répliques générées, ces techniques peuvent offrir au modèle une opportunité de se généraliser. Cependant si trop d'images sont générées, trop peu modifiées ou dénaturées, l'augmentation de données desservira l'apprentissage menant à une non convergence ou au contraire à un over-fitting du modèle. C'est pourquoi ces modalités d'augmentation de données peuvent être incluses dans les paramètres du sweep pour obtenir un compromis favorable.

iv. Les modèles

Aujourd'hui, il existe de nombreux modèles de détection d'objets en libre accès, dont la plupart sont pré-entraînés. Cela signifie que des paramètres sont déjà attribués pour un certain nombre de classes d'objets, souvent assez génériques, pouvant identifier les objets du quotidien : table, chaise, voiture, maison etc. Pour ce type de modèle, on peut citer **GroundingDino** (IDEA Research), entraîné sur un lexique d'environ 30 000 mots, organisés en unités de bases (racines, préfixes, suffixes ...) afin de reconnaître un plus grand nombre de mots par combinaisons (Devlin et al. 2019). De plus, le prompt

(i.e. instruction d'entrée du modèle) peut se composer d'une combinaison sémantique de mots, ce qui multiplie les possibilités de détection. En sortie, le modèle génère des boîtes englobantes autour des objets concernés par le prompt. Chaque détection est associée à un score de confiance du modèle en sa prédiction, à interpréter relativement à la détection d'un autre objet de la même classe sur l'image afin de comparer la certitude accordée. Ce score de confiance, compris entre 0 et 1, ne reflète pas la probabilité que le modèle se trompe pour la détection donnée, mais est calculé relativement aux autres sorties du réseau de neurones, pour lesquels le modèle accorde moins de confiance.

Associé à ce modèle de détection, **SAM, Segment Anything Model**, développé par le groupe Meta, permet de segmenter l'image au sein de la boîte englobante (Grounded-SAM Contributors 2023). L'association Grounded-SAM est bien documentée en ligne, notamment sur la plateforme Github ou directement sur le site des concepteurs. La segmentation est une opération identifiant les pixels faisant partie d'une entité visuelle cohérente commune, au travers d'une description de l'état de chaque pixel par une matrice indicatrice (i.e. masque) remplie de 0 ou 1, booléens ou 0 et 255. Ce modèle prend comme entrée une image et un prompt sous forme d'une boîte englobante ou d'un point. Selon la nature du prompt et de sa position sur l'image, un masque est généré pour segmenter la forme ciblée.

Pour autant, Grounding Dino n'est pas en mesure de détecter des classes d'objets très spécifiques sur les images, malgré des tentatives d'utiliser un lexique descriptif des objets cibles. C'est pourquoi, nous avons recours à des modèles de classification et détection d'objets, très répandus pour réaliser des tâches spécifiques. Les modèles Yolo de chez *Ultralytics* sont de parfaits candidats, performants et très documentés (Jocher, Qiu, Chaurasia 2023). Ces modèles sont également pré-entraînés sur un grand jeu de données tel que *COCO* (Common Object in COntext) afin de simplifier la phase d'entraînement sur de nouvelles classes, car les paramètres ont déjà été pré-estimés pour la détection d'autres objets. Cette phase appelée fine-tuning consiste à s'appuyer sur le pré-entraînement du modèle pour apprendre de nouvelles classes, basées sur une labellisation spécifique, propre au projet.

Depuis la sortie en 2015 de Yolo, *Ultralytics* n'a cessé d'améliorer leur modèle jusqu'à la version la plus récente **Yolov11**. Offrant de meilleures performances que ses précédentes versions (cf. Figure 3), la version se décline en différents modèles (n, s, m, l et x) dont la taille et la complexité varient. Plus le modèle est complexe, plus il sera efficace à extraire des propriétés visuelles des objets étudiés, mais mobilisera plus de ressources informatiques ou prendra davantage de temps à s'exécuter. Ces modèles et leurs packages offrent des outils d'entraînement et d'évaluation des performances intuitifs et efficaces, car sont basés sur des modèles pré-entraînés, sur lesquels l'apprentissage est plus performant que de partir « from scratch » (Shallu, Mehra 2018).

Il existe de nombreux autres modèles d'analyse d'images en apprentissage profond pour de la détection et la segmentation, axant leur stratégie sur l'instantanéité ou sur d'autres performances (Arkin et al. 2023).

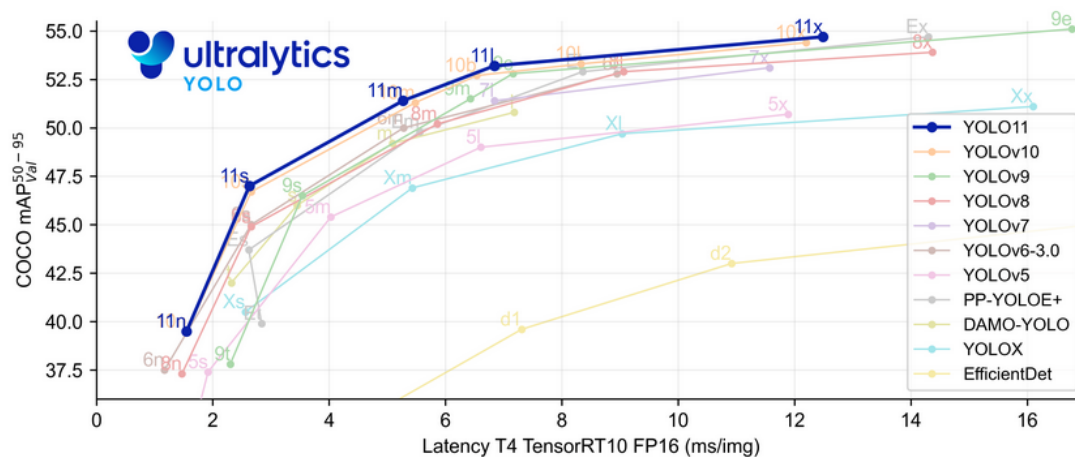


Figure 3 : Performances des modèles Yolo

Matériel et méthodes

1. Données

a. Base de données

Notre étude se base sur des photographies prises par les opérateurs d'une station de piégeage située sur la Bresle (Hauts de France et Normandie) à l'aide de leur smartphone, permettant l'acquisition d'images de résolution 4000*3000 pixels. Tous les poissons sont capturés au cours de leur migration anadrome, lors de laquelle les espèces diadromes remontent les cours d'eau vers les zones de frayères pour la reproduction, et sont photographiés systématiquement sur les deux flans. D'autres prises de vues annexes peuvent être ajoutées, zoomant sur une zone particulièrement impactée par une pathologie. Sur chaque photo, une étiquette de référence de 3 cm, plastifiée et réutilisée sur chaque image, est placée dans le champ de la prise de vue et le poisson est posé à plat sur une table blanche (cf. Figure 4). L'échelle standardisée nous permet d'effectuer une conversion cm/pixel pour connaître la taille des individus et évaluer la surface de recouvrement des dommages. Au total, 335 individus différents ont été photographiés correspondant à 1039 images pour l'année 2024 et le début de la saison 2025 (cf. Tableau 2).



Figure 4 : Exemple d'image acquise par les opérateurs

Tableau 2 : Nombre d'individus disponibles pour notre étude par année et par espèce

Espèce/Année	2024	2025
Truite de mer	261	69
Saumon atlantique	2	2
Truite fario	1	0
Total	264	71

Ces images alimenteront l'apprentissage du modèle Yolo de détection des pathologies, qui va extraire les caractéristiques des objets cibles mais également les informations de contexte de l'image. Cela permet de mieux détecter en identifiant des patterns récurrents. Le modèle va donc s'adapter aux informations visuelles de contexte autour des objets et risque de perdre en efficacité dès lors que l'hétérogénéité de l'arrière-plan changera. C'est pourquoi, si l'on souhaite pouvoir appliquer notre modèle ultérieurement à des images moins uniformisées, qu'une table blanche propre, il semble nécessaire de s'affranchir de ce contexte en ne conservant que l'image centrée sur le poisson. Pour ce faire, le modèle Grounding Dino est utilisé pour détecter le poisson et disposer une boîte englobante autour du corps de l'individu. La portion extérieure de la boîte est ensuite rognée. De cette manière, un poisson pris sur l'herbe ou une table de biométrie a plus de chance d'être bien analysée malgré un contexte différent. Ce sont donc ces images centrées sur le poisson qui seront utilisées tout au long de l'étude.

Chaque poisson est identifié par un identifiant unique à l'individu, contenu dans le nom du fichier de l'image. Cette information permet de relier l'animal présent sur la photographie à ses caractéristiques bancarisées, notamment sa biométrie (taille, poids) et les codes pathologie définis par l'opérateur. Les pathologies renseignées sur le poisson ne sont pas nécessairement visibles sur chaque image de l'individu. C'est pourquoi des prises de vues exhaustives sont impératives pour visualiser l'ensemble des dommages de l'individu étudié. Ainsi, les photos des deux flancs du poisson nous apporteront le maximum d'information, et les autres zooms compléteront les prises de vues générales, notamment pour les poux, très présents autour de l'anus, ou sur le dos.

Dans une démarche exploratoire des méthodes d'analyse d'images, nous avons restreint notre étude à trois groupes de pathologies (pertes d'écailles, rougeurs et poux de mer) et sept localisations, agrégations de plusieurs unités de localisations précises de la littérature. En effet, les codes patho disponibles sur les images de notre jeu de données issue de la Bresle contiennent quelques pathologies largement majoritaires (cf. Tableau 3). Les absences d'organes (AO), troisième pathologie la plus représentée, ont été omises de l'étude car la méthode de détection d'absence d'un objet diffère de la détection de présence pour notre modèle.

De même, l'agrégation des localisations a été basée sur la récurrence de chaque code (cf. Tableau 4) et sa simplicité à être identifiée automatiquement sur l'image. La bouche (MT) que l'on retrouve figure 1, septième localisation la plus fréquente, est comprise dans l'unité morphologique de la tête (TT). Ici, le « T » en deuxième lettre est un suffixe faisant référence à la zone indiquée, de façon totale et non une sous-unité. Cette simplification permet de s'affranchir de la complexité et

l'exhaustivité des codes théoriques en regroupant de nombreuses sous-unités en agrégats comme la tête, le haut et le bas du corps etc.

Tableau 3 : Extrait des codes relatifs aux pathologies

	Code	Nombre	Proportion (%)	Code agrégé
Perte d'écailles	EC	1574	28.5	EC
Erosion d'écailles	ER	1243	22.5	EC
Poux (avec et/ou sans flagelles)	PU	386	7.0	PU
Rougeur	RO	378	6.8	RO
Pétéchie	PE	320	5.8	RO
Présence d'une ou plusieurs lésions à la surface du tégument	PL	241	4.4	RO

Tableau 4 : Extrait des codes relatifs aux localisations des pathologies

	Code	Nombre	Pourcentage (%)
Tronc	WT	1688	30.6
Tête	TT	713	12.9
Tout le corps	CT	702	12.7
Nageoire caudale	QT	628	11.4
Pédoncule caudal	KT	602	10.9
Abdomen	AT	552	10
Dos / haut du corps	HT	463	8.4

Le nombre de photographies par mois dépend de la phénologie des espèces. Ainsi, près de deux tiers des images disponibles proviennent des mois de juin et juillet. Les poissons étant à cette période en phase migratoire, ils arborent une robe argentée, homogène jusqu'au mois d'août. C'est à partir du mois d'octobre que les individus changent pour atteindre leur robe de reproduction, plus colorée et tachetée. Par ailleurs, c'est au cours de la période estivale que l'on rencontre la plus grande proportion de poissons observés atteints par les pathologies étudiées sur l'année 2024 (cf. Figure 5). Par exemple, les poux de mers se décrochent du corps des salmonidés au cours de leur séjour en eau douce (Heffernan 1999). Cela explique la plus forte occurrence des poux l'été, durant la montaison, et moins en hiver, quand les poissons ont séjourné plusieurs mois en milieu dulçaquicole. Ainsi, les données disponibles utilisées lors de cette étude, issues de la saison estivale, contiennent une forte représentativité des pathologies d'intérêts.

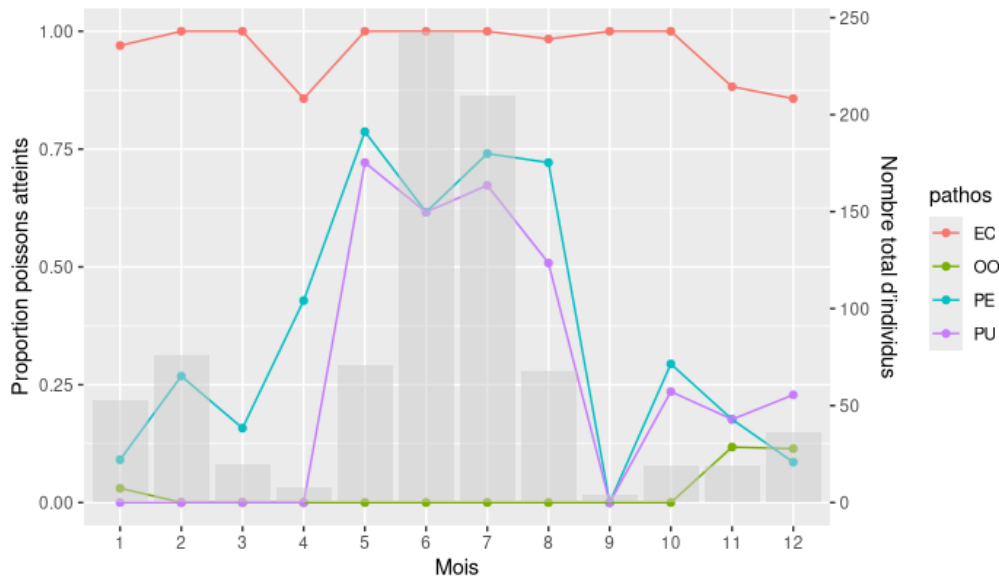


Figure 5 : Saisonnalité de l'occurrence des pathologies et nombre de poissons rencontrés

2. Modèles

Grounding Dino est utilisé pour identifier sur l'ensemble des images l'étiquette d'échelle et le poisson quand il est entier ou encore pour détecter la tête, l'œil et la nageoire caudale des individus. Ainsi, pour l'identification de ces objets, ce modèle est favorisé, car clé en main, sans entraînement spécifique supplémentaire pour la détection et nous permet de nous affranchir de la phase pré-opérationnelle chronophage de labellisation et d'entraînement du modèle.

D'autre part, SAM est utilisé dans le cadre de ce projet pour extraire le masque du poisson, pour connaître sa surface et délimiter finement sa forme et ses contours. De plus, ce modèle pourra affiner la surface occupée par une pathologie après détection par une boîte englobante. A l'instar de Grounding Dino, le modèle SAM est rapidement opérationnel et compatible avec la nature des détections des autres modèles utilisés.

Enfin, les modèles Yolo sont utilisés pour la détection des pathologies sur le corps des poissons, problématique très spécifique à notre projet. Ces modèles sont des outils pour lesquels la phase d'apprentissage est très documentée. Nous avons sélectionné le modèle Yolo11m, dont le compromis entre le gain de performance et l'augmentation du temps de calcul est raisonnable pour nos ressources informatiques (cf. Figure 3).

3. Construction du jeu de données et apprentissage

a. Labellisation

Pour permettre au modèle Yolo11m, pré-entraîné sur des classes usuelles, de détecter les pathologies étudiées, une phase de labellisation est mise en place. Cette étape permet de fournir au réseau de neurones un jeu de données d'entraînement renseigné par des opérateurs sur laquelle le modèle va baser son apprentissage.

Pour ce stage, cette tâche a été réalisée au travers d'une interface Shiny, développée lors d'un précédent stage dans l'unité (Dubois 2024). Cette application en ligne est reliée à une base de données pour stocker les informations entrées. J'ai adapté l'interface (Rshiny) pour répondre aux problématiques spécifiques du projet. Entre autres, nous avons imposé une lecture des images pathologie par pathologie afin de concentrer l'effort d'analyse du labellisateur sur une seule classe à la fois.

Cette phase de labellisation nécessite une certaine précision de la part du labellisateur pour renseigner de façon homogène toutes les informations, en prêtant attention aux détails de l'image. C'est pour nous affranchir de ce processus chronophage d'acquisition des labels, dont les performances de l'entraînement sera sensible, que nous avons préféré nous appuyer sur les performances de Grounding Dino pour la détection du poisson, la tête et l'œil. En effet, selon le labellisateur et sa méthode d'identification des objets, l'apprentissage du modèle peut différer, si la détection n'est pas une information très explicite (Rädsch et al. 2023).

Les données entrées sur l'application Shiny alimentent directement une base de données sur Postgres, contenant toutes les informations de l'image et la labellisation effectuée. Cela permet de stocker et mettre à jour les données, ainsi que d'y avoir accès par téléchargement pour l'entraînement. La structure de cette base a été modifiée par rapport à sa version originale, pour visualiser les pathologies de notre étude et répondre à la nouvelle structure de l'application.

Par soucis de temps, une première étape a été réalisée par mes encadrants et moi-même, avec un regard « naïf » sur la caractérisation des pathologies. Puis, les opérateurs, experts sur le sujet de l'identification des pathologies, ont été mobilisés pour relire et ajuster cette labellisation. Ils se sont placés en tant que correcteurs sur l'application pour avoir accès à nos boîtes englobantes et modifier directement les informations sur la base de données qui constituera le jeu de données utilisé lors de la phase d'apprentissage du modèle Yolo.

b. Paramètres d'apprentissage

La phase d'apprentissage du modèle lui permet d'apprendre la détection de nouvelles classes d'objets. Il faut fournir au modèle un jeu de données d'entraînement, de validation et de test (cf. 2.b) comprenant les images, les boîtes englobantes et la classe de l'objet associé. C'est sur ces données que l'entraînement sera lancé dans le cadre d'un algorithme de balayage (ou sweep), permis par Weight and Biases (W&B 2025), afin d'extraire la meilleure combinaison d'hyperparamètres qui sera définie pour l'apprentissage du modèle définitif.

c. Augmentation de données

Avec un jeu de données photographiques inférieur à mille images, l'occurrence de chaque pathologie est limitée, il est donc nécessaire d'utiliser une méthode d'augmentation de données. Les changements d'échelles de l'image ont été omis lors de la phase d'entraînement du modèle de détection des nageoires caudales, car l'objet est situé à l'extrémité de l'image droite ou gauche et a un fort risque d'être rogné ou supprimé, ce qui perturberait le modèle. La modification de la teinte, la saturation et la luminosité, la translation, le changement d'échelle et les rotations sont les opérations appliquées aux images du jeu de données. En s'appuyant sur cette méthode, nous avons généré deux nouvelles images à partir de chaque image brute de notre jeu de données.

Tableau 5 : Nombre d'occurrence des pathologies utilisées lors de l'apprentissage de Yolo

Pathologies	Entraînement	Validation	Test	Total
EC	1600	207	145	1952
RO	600	78	47	725
PU	550	75	43	668

C'est donc à partir de ces données augmentées et des hyperparamètres identifiés par le sweep comme optimaux, que l'entraînement du modèle va pouvoir s'effectuer.

4. Estimation de la gravité

A l'issue de cette phase d'apprentissage, les pathologies détectées par le modèle sont identifiées par une boîte englobante, une classe et un score de confiance associé. Nous pouvons estimer, à partir de cette détection, l'importance de chaque maladie sur le corps du poisson en quantifiant sa surface relative sur le corps du poisson (EC, RO) ou en comptant le nombre de parasites qui y sont fixés (PU).

Pour obtenir ces codes de gravité par pathologie, par individu, il faut étudier et compiler les informations des photographies des deux flancs du poisson. Il suffit de comptabiliser le nombre de poux présents (PU) pour obtenir la classe de gravité de cette pathologie alors que nous devons extraire la surface totale concernée par les autres pathologies externes (EC, RO). Pour ce faire, le modèle SAM segmente l'objet présent dans les boîtes englobantes des détections, puis un masque cumulatif par classe d'objet est créé sur chaque image et la surface en pixels de ce masque est calculée. Passer par la segmentation de chaque détection permet d'affiner l'estimation de la surface occupée par la pathologie d'une façon plus précise que calculer la surface de la boîte englobante, encadrant grossièrement l'objet. Après compilation des résultats pour les des deux faces, nous sommes l'occupation de la pathologie et divisons par la surface du masque du poisson, calculée suite à l'utilisation du modèle SAM, pour obtenir un pourcentage de recouvrement, et en déduire la classe de gravité, conformément aux indices usuels (cf. Tableau 1).

5. Estimation de la localisation

a. Partition du poisson

Par la suite, nous devons localiser chaque détection de pathologie sur le corps du poisson. En effet, le modèle est naïf sur la/les parties du corps concernées par chaque pathologie. Il y a donc une phase de post-traitement qui intervient pour obtenir la partie localisation du code patho.

Chaque unité morphologique est délimitée à l'aide d'une stratégie adaptée, afin d'être le plus généralisable possible et s'ajuster au cas du saumon et truite de mer, proches morphologiquement. Le travail de localisation consiste à transcrire de manière automatique la partition, évoquée sur la figure 1, sur les images analysées.

b. Courbure et longueur du poisson

Les séparations de chaque unité morphologique sont définies comme orthogonales à l'axe central du poisson, et dans certains cas, ces délimitations sont placées au niveau d'un ratio de la longueur totale du poisson. C'est pourquoi l'obtention de la longueur du poisson et de sa courbure sont des points essentiels pour atteindre nos objectifs.

Dans un premier temps le poisson est identifié par Grounding Dino et est segmenté par le modèle SAM pour pouvoir éliminer toutes les détections situées hors du masque du poisson. Dans la plupart des situations, le corps des individus est courbé sur la prise de vue, ce qui va influencer sur la séparation des différentes unités morphologiques du poisson. Pour extraire la courbure du poisson, on applique un algorithme de squelettisation du masque du poisson qui est une méthode itérative consistant au retrait successif de couches de pixels pour ne conserver que les pixels médians de l'image, par érosion. En ne conservant que le plus grand segment de ce squelette, on obtient une approximation de la courbure du poisson, décrite par un ensemble de coordonnées de pixels constituant le squelette (cf. Annexe I). Le tracé de ce segment remonte jusqu'à l'extrémité la plus éloignée de la nageoire caudale, pouvant biaiser l'estimation de cette courbure. A l'opposé du corps, l'angle produit par l'ouverture de la bouche va également influencer sur le tracé du squelette. Le tracé est donc réduit au segment entre le masque de la tête et celui de la nageoire caudale. En utilisant ces masques et la géométrie de unités morphologiques, nous avons extrait les coordonnées de l'extrémité de la bouche et le point fourche de la nageoire caudale (cf. Annexe II). Ces deux points nous permettent d'effectuer une régression de la courbure du squelette par un polynôme du second degré, dont le tracé est forcé de passer par les points définis (cf. Annexe III).

Dans la suite de la démarche, nous allons utiliser un ratio de cette longueur pour placer des segments d'intérêt. Nous devons donc nous assurer que la taille prédite par la méthode ci-dessus est proche de la longueur réelle du poisson. Pour évaluer l'efficacité de cette approximation, nous pouvons nous baser sur la surface du masque de l'étiquette de 3 centimètres (Détection par Grounding Dino, puis segmentation par SAM) comme facteur de conversion pixel – centimètre. Différentes méthodes sont comparées pour minimiser l'écart entre la taille mesurée par l'opérateur et la taille estimée à partir de la photographie.

c. Identification des unités morphologiques

Une fois l'axe central du poisson obtenu, la délimitation de la tête est définie comme la droite perpendiculaire à la courbure, dans le masque du poisson, juste à l'extrémité du masque de la tête. Le même raisonnement est suivi pour la délimitation de la nageoire caudale (cf. Figure 1).

La délimitation corps/queue est elle aussi perpendiculaire à la ligne de courbure de l'individu et définie comme un ratio de la longueur totale. En effet, la position de l'anus se situe à un pourcentage de la longueur du poisson. D'après les valeurs bancarisées par FishBase (FishBase 2025a; 2025b), la limite anale est positionnée à 69 % de la longueur totale. Ainsi, la délimitation corps/queue est placée à x_{lim} , tel que :

$$x_{lim} = \alpha \times Length_{total} = \alpha \times \beta \times Length_{fork}$$

Où :

- α : pourcentage de la longueur totale de la position de l'anus
- β : ratio entre la longueur fourche et la longueur totale
- On a $\alpha * \beta = 0.69$ (FishBase 2025a)

La même approximation est faite pour les deux espèces d'études sans distinction, bien que le pédoncule caudal soit plus long chez le saumon que la truite de mer (FishBase 2025a).

Enfin, la position de la ligne latérale est importante pour distinguer les parties supérieure et inférieure du corps. Pour cela, nous reprenons l'équation de la courbure du poisson mais forçons le tracé à de nouveaux points de référence afin d'opérer à une translation de celle-ci. Pour ce faire, il est nécessaire d'obtenir une relation de proportionnalité entre la hauteur des extrémités de la ligne latérale par rapport à la hauteur du poisson à ces niveaux, afin d'avoir un ratio de hauteur à appliquer à l'axe central du découpage de l'individu. Nous avons donc basé cette exploration sur les mesures de vingt poissons.

Toutes ces opérations de ratio aux différentes longueurs sont orientées pour une orientation précise du poisson. Il est donc nécessaire de tenir compte du sens (i.e. face ventrale en bas ou en haut de l'image) et de l'orientation (i.e. tête à gauche ou à droite de l'image), représentés figure 9. La recherche de l'orientation consiste à comparer les ordonnées (axe vertical) des centres des boîtes englobantes de la tête et la nageoire caudale. Tandis que le sens est déterminé en fonction de la position du centre de la boîte de la tête par rapport au squelette du masque de la tête, reflétant l'axe médian du poisson. Ainsi, le sens et l'orientation sont décrits par la valeur 1 si le poisson est face ventrale vers le bas et tourné vers la droite, -1 sinon. De plus, pour faciliter le raisonnement et la régression du squelette du poisson par un polynôme de degré 2, les poissons sont identifiés sur l'image par Grounding Dino, puis l'extérieur de la boîte englobante est rogné et l'image est horizontalisée.

Une fois ces limites de localisation placées, nous pouvons comparer le recouvrement de chaque détection d'une pathologie au découpage et identifier les unités morphologiques atteintes. Comme vu précédemment, il existe des codes agrégés de localisation. C'est pourquoi, le choix du code final de sortie est défini par un arbre de décision complexe en fonction des combinaisons de codes (cf. Annexe IV). Cet arbre n'est pas hiérarchique et est structuré comme une succession de conditions binaires d'absence ou présence de la pathologie sur chaque unité morphologique, aboutissant à un code, parfois agrégé, parfois détaillé.

6. Evaluation des performances du pipeline

a. Métriques usuelles du modèle de détection

En analyse d'images, par apprentissage profond, il existe de nombreuses métriques d'évaluation des performances de détection du modèle. Parmi elles :

- Precision (précision) : proportion de détections qui sont réellement correctes, soit $P = \frac{TP}{TP+FP}$, avec TP le nombre de vrais positifs et FP le nombre de faux positifs. Ainsi cette métrique tend vers 1 avec la capacité du modèle à détecter les bons objets recherchés.

- Recall (rappel) : proportion d'objets réels bien détectés par le modèle, soit $R = \frac{TP}{TP+FN}$, avec TP le nombre de vrais positifs et FN le nombre de faux négatifs. Ainsi cette métrique tend vers 1 avec la capacité du modèle à être exhaustif dans sa détection.
- mAP50 (mean average precision at IoU 50 %) : performance globale prenant en compte la précision et le rappel pour les détections ayant un IoU (intersection over union) d'au moins 0.50, avec $IoU = \frac{Aire_{intersection}}{Aire_{union}}$, (Lin et al. 2015)
- mAP50-95 (mean average precision at IoU from 50 to 95 %) : Moyenne des mAP pour dix seuils d'IoU de 50 à 95 %, soit $mAP_{50:95} = \frac{1}{10} \sum_{IoU=0.5}^{0.95} mAP_{IoU}$. Cette métrique résume le rappel et la précision tout en pénalisant davantage les imprécisions de localisation des détections.

Ces métriques permettent d'évaluer les performances du modèle après sa phase d'apprentissage et de discriminer un modèle d'un autre mais ne reflètent pas les réelles capacités à détecter les pathologies. En effet, la définition de ces objets est sujette à la subjectivité du labellisateur ou de l'opérateur sur le terrain. L'humain, référence de ce travail, ne pénalise pas autant le manque de précision que les métriques de performances automatiquement générées. Par exemple, le taux de recouvrement entre la boîte prédite et la boîte labellisée peut être faible, mais valider tout de même la détection. Une boîte englobant une grande tâche diffuse, définie par le labellisateur peut être détectée en plusieurs unités par le modèle, qui va pénaliser sa détection, à cause du manque de similarité avec sa référence. Pour savoir si le modèle détecte bien les pathologies dans l'ensemble et pouvoir comparer deux modèles différents, on peut baser l'évaluation sur l'écart entre l'indice de gravité estimé et l'indice renseigné dans la base données et comparer indépendamment les localisations par la suite.

b. Prédiction de la gravité

Comme évoqué plus haut, le modèle évalue sa détection et classification grâce à un score de confiance du modèle en sa détection, associé à chaque boîte englobante. Pour comprendre l'influence du seuil de confiance accordé pour valider ou non les détections du modèle, nous avons réalisé une évaluation des performances du modèle à prédire le bon indice de gravité. Nous avons donc sélectionné dix seuils (0.01, 0.03, 0.05, 0.1, 0.2, 0.3, 0.5, 0.7, 0.8, 0.9) de confiance, en deçà duquel les détections ne sont pas considérées comme valides, et calculé la différence entre la valeur observée et la valeur d'indice de gravité estimée. Cela nous permet d'obtenir des probabilités discrètes des erreurs d'estimation de l'indice à différents seuils de confiance. Afin de tenir compte de l'incertitude concernant l'estimation de l'indice de gravité par les opérateurs (cf. Tableau 1), nous avons poursuivi le raisonnement d'évaluation en considérant qu'un écart d'une unité de l'indice était une bonne prédiction. En effet, un opérateur ne mesure pas mais évalue visuellement la surface occupée par la pathologie. Cette évaluation reflètera davantage les capacités du modèle avec comme référence l'expertise de l'opérateur. Par construction, un delta nul entre la prédiction et l'observation reflète un indice de gravité conforme à ce que l'expert a observé.

c. Prédiction du code de localisation

De même, une évaluation du modèle sur sa capacité prédictive du code de localisation permettra de compléter l'analyse des performances de détection. En effet, un indice de gravité correct ne nous renseigne pas si le modèle sélectionne les bons objets sur l'image. Il nous informe juste sur le fait que la surface d'objet estimée par le modèle est conforme à la classe de gravité notifiée par l'opérateur. En ajoutant l'information de la localisation, on affine notre comparaison des estimations aux codes observés.

Pour ce faire, plusieurs méthodes ont été envisagées. L'une consistait à comparer l'exactitude du code fourni par rapport au code observé. Cependant, discriminer sévèrement un code généré alors qu'il recoupe potentiellement le code fourni les opérateurs est une méthode qui ne tient pas compte de la subjectivité apportée lors de l'observation. C'est pourquoi nous avons cherché à quantifier la distance entre la prédiction et l'observation, ce qui ne peut être mis en place car l'arbre de décision établi (cf. Annexe IV) n'est pas hiérarchisé mais arbitraire. De plus, au vu du nombre de cas particuliers et combinaisons possibles, définir une note à chacune d'entre elles semble aussi fastidieux que la notation manuelle. C'est pourquoi, nous avons, manuellement, attribué une note entre 0 et 1 de chaque prédiction du jeu de données test, afin d'établir une fonction d'utilité (Von Neumann, Morgenstern 1944) :

- 0 : le code ne correspond pas à l'observation
- 0.33 : la surface de la pathologie prédite est plus large que ce qu'a observé l'opérateur (surestimation)
- 0.66 : la surface de la pathologie prédite est moins étendue que ce qu'a observé l'opérateur (sous-estimation)
- 1 : le code issu de la prédiction est identique au code défini par l'opérateur

L'application d'un principe de précaution nous a amené à favoriser une prédiction plus stricte qu'une prédiction trop large de l'étendue de la pathologie. Ainsi l'espérance d'utilité empirique par seuil de confiance est calculé.

d. Performances du modèle sélectionné

Ce travail d'évaluation des performances prédictives de la gravité et localisation nous permettra de d'identifier, par pathologie, un ou plusieurs seuils de confiance pour obtenir le meilleur résultat prédictif. En effet, pour une pathologie donnée, un seuil peut maximiser le score de gravité, tandis qu'un autre correspond à la meilleure performance du modèle à prédire la localisation en moyenne.

A l'issue de ce choix, nous pourrions quantifier les performances du modèle de détection au regard de sa capacité à prédire une absence ou une présence d'une pathologie sur le poisson. Plusieurs indicateurs résument la matrice de confusion entre les observations et les prédictions et nous apportent des informations spécifiques pour l'évaluation :

- Précision : capacité du modèle à bien prédire, absence et présence, calculé comme le quotient : $P = \frac{TP+TN}{TP+FP+TN+FN}$, en se basant sur la matrice de confusion. On retrouve cette métrique dans l'évaluation usuelle des performances du modèle de détection, avec des critères

différents (non basé sur de l'absence/présence, mais sur une invalidation des détections selon un critère de recouvrement)

- Sensibilité : reflète la capacité du modèle à bien détecter les positifs observés, calculés comme : $\text{Sens} = \frac{TP}{TP+FN}$
- Spécificité : correspond à la capacité du modèle à prédire correctement les négatifs observés, calculé comme : $\text{Spe} = \frac{TN}{TN+FP}$
- Kappa de Cohen : mesure l'accord entre les observations et prédictions. Il permet quantitativement de comparer les performances prédictives du modèle et le hasard. Il est le résultat de l'opération $k = \frac{P_A - P_H}{1 - P_H}$, avec :
 - P_A : proportion d'accord entre l'observation et la prédiction, soit la précision du modèle.
 - P_H : probabilité d'accord avec le hasard entre observation et prédictions du modèle, décomposée comme la somme d'une probabilité d'accorder des positifs et d'accorder des négatifs : $P_H = \frac{(TP+FP)(TP+FN)}{TP+FP+TN+FN} + \frac{(FN+TN)(FP+TN)}{TP+FP+TN+FN}$.

Plus k est grand, plus l'accord entre observation et prédiction est important. Au contraire, si k est faible, la part d'aléatoire est grande. De plus, par la nature du calcul, si $k < 0$, alors on conclut à un désaccord et à une part plus importante au hasard qu'aux performances du modèle à générer de bonnes prédictions.

7. Application à un cas concret

Une fois notre méthode mise en place, nous pouvons appliquer l'ensemble du pipeline (cf. Annexe V) sur un ensemble d'individus de notre jeu de données issue de 2025, pour lequel le modèle de détection est naïf afin de prédire les pathologies présentes. Comme nous disposons des informations fournies par les opérateurs, nous pourrions comparer le résultat prédit aux observations sur le terrain.

Certains parasites, tels que les poux de mer (*Lepeophtheirus salmonis*), ont un taux de survie très faible en eau douce et finissent par se décrocher du corps des salmonidés (Heffernan 1999). Ce sont donc de bons indicateurs du temps de séjour des poissons en milieu dulçaquicole. Nous pouvons donc évaluer la présence de poux sur les individus étudiés pour savoir s'ils sont arrivés récemment en eau douce, ou s'ils y séjournent depuis longtemps, et ainsi suivre l'évolution temporelle de ce parasitisme. Cette étude s'étend du mois d'avril au début du mois de juillet 2025, dates pour lesquelles nous avons des images issues de la station de la Bresle.

8. Déroulé du pipeline

Après avoir détaillé la démarche suivie étape par étape de la méthode d'obtention d'un code patho à partir d'une image, nous pouvons récapituler le pipeline (Figure 6) :

- L'image prise par les opérateurs est analysée par Grounding Dino pour détecter le poisson (1) et rogner le contexte autour de l'objet (2).

- Le modèle Yolo entraîné pour détecter les pathologies analyse l'image et renseigne les boîtes englobantes et la classe de chaque détection (3).
- SAM est utilisé pour segmenter la forme de la pathologie au sein de la détection afin d'affiner la surface concernée (4).
- Le dénombrement des parasites ou la somme cumulée de la surface d'une pathologie est étudié afin d'obtenir un code de gravité selon les critères usuels (5).
- En parallèle, l'identification des unités morphologiques du poisson, à partir de l'image vierge, permet d'obtenir la localisation de chaque détection, par la suite compilées selon l'arbre de décision définit plus haut (6).
- Ces informations sont donc compilées par pathologies afin d'obtenir un code patho prédit (7).

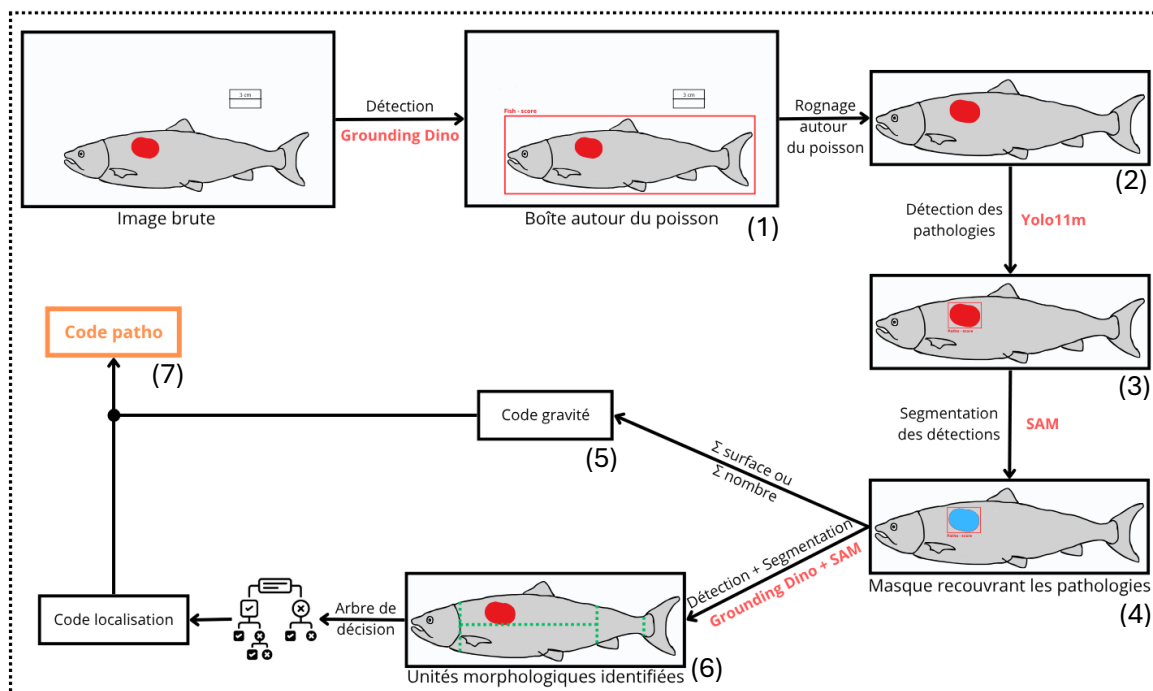


Figure 6: Extrait du schéma du pipeline (b. en annexe V) d'extraction des codes pathos d'un individu

Résultats

1. Opérationnalité du découpage morphologique

Plusieurs étapes de l'obtention des unités morphologiques sont basées sur les performances de modèles pré-entraînés et impliquent l'estimation de plusieurs variables. Afin de s'assurer de la rigueur de notre méthodologie, nous devons évaluer ces étapes intermédiaires permettant de prédire la localisation de chaque pathologie.

Premièrement, l'utilisation de modèles pré-entraînés comme Grounding Dino, peut comporter des lacunes dans la détection de certains objets très spécifiques. C'est le cas de la nageoire caudale, qui, sur près d'un tiers des images, n'est soit pas détectée, soit confondue avec la nageoire dorsale ou le poisson complet, car « nageoire caudale » est sémantiquement proche de « poisson ». C'est pourquoi, nous avons également entraîné un modèle Yolo spécifique à la détection de la nageoire caudale pour

palier à la perte de ces images qui se seraient retrouvées inexploitable pour extraire la localisation de la maladie. Malgré cette correction, sur certaines images à l'éclairage ou cadrage différent, notre modèle Yolo de détection de la nageoire caudale ou Grounding Dino pour l'œil manquent certains objets. Ainsi, 5 % des images (soit 16 images) ne sont pas exploitables et n'entrent pas dans notre phase d'évaluation des performances du pipeline.

Par ailleurs, l'estimation de la courbure du poisson par la régression du tracé du squelette permet d'estimer la longueur fourche de l'individu. Comme certaines limites d'unités morphologiques sont définies comme un ratio de cette longueur, l'erreur d'estimation de celle-ci est quantifiée. En comparant les valeurs obtenues à celles renseignées sur la base de données, nous obtenons une erreur absolue d'estimation moyenne de 1.9 % de la longueur du poisson (cf. Figure 7), meilleur résultat d'estimation correspondant à l'utilisation de la courbure en se basant sur le rectangle de 3 cm comme échelle. Lorsque que l'on compare les mesures de longueurs par un logiciel comme *ImageJ* aux données acquises sur le terrain, on remarque que les valeurs prédites par l'algorithme sont plus proches de la précision d'*ImageJ* que les mesures sur le terrain.

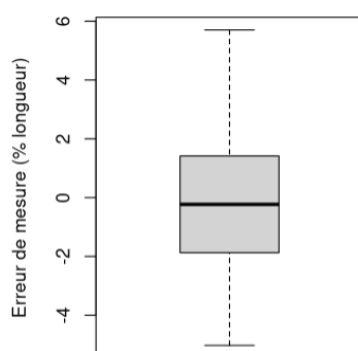


Figure 7 : Distribution des erreurs de mesures de la longueur fourche (% de la longueur de référence)

Après l'obtention du tracé de la courbure, nous appliquons, selon la méthode, un facteur correctif de la hauteur par laquelle doit passer la ligne latérale. En étudiant sur 20 poissons, le ratio entre la hauteur de la ligne latérale et la hauteur de la face latérale du poisson, on remarque une relation affine (cf. Figure 8 a, b). Il existe ainsi une proportionnalité entre la hauteur du poisson et la ligne latérale, de 70 % au niveau de la tête et 56 % proche de la queue. Après application de cette correction, nous traçons la ligne latérale entre ces points.

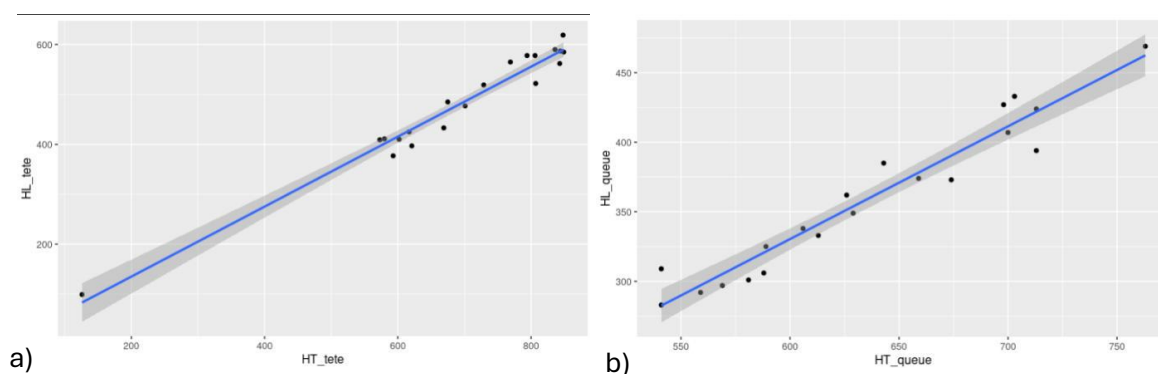


Figure 8 : Relations entre la hauteur totale du poisson au niveau de la tête et la hauteur de la ligne latérale (a.) et entre la hauteur totale du poisson au niveau du pédoncule caudal et la hauteur de la ligne latérale (b.)

Enfin, le tracé automatique de chaque unité morphologique, suivant la méthodologie est effectué (cf. Figure 9), afin d'identifier la tête, le corps du poisson avec les zones supérieure et inférieure, la queue et la nageoire caudale du poisson. La séparation du corps s'ajuste correctement à la courbure de la ligne latérale de l'individu.

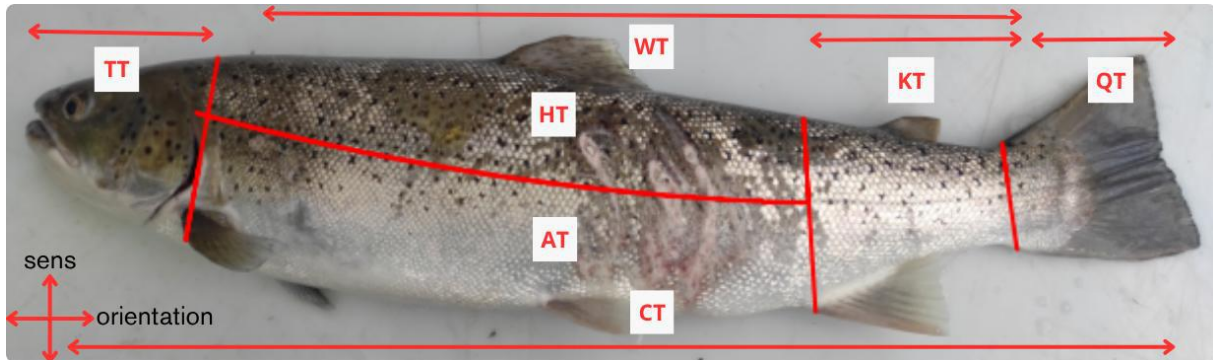


Figure 9 : Partitionnement morphologique automatisé du poisson

2. Evaluation de l'apprentissage du modèle de détection des pathologies

a. Sélection du meilleur modèle

Lors du sweep, l'algorithme va explorer l'espace des hyperparamètre afin d'obtenir le modèle maximisant les métriques cibles de performances. Une fois le sweep achevé, nous obtenons une comparaison de modèles ayant des combinaisons d'hyperparamètres différents. Nous pouvons les discriminer selon différentes métriques, que l'algorithme renvoie à chaque fin d'apprentissage d'un modèle.

La valeur de la perte associée à la prédiction des boîtes englobantes (i.e. *box loss*) à l'issue de l'entraînement est la métrique cible, que le réseau de neurones visait à minimiser. Un palier minimum est atteint dès le cinquième modèle (annexe VI). Le nuage de points dispersés, correspond à l'exploration des combinaisons d'hyperparamètres ayant entraînés des modèles dont les performances s'éloignent de l'optimum trouvé. On peut donc dans un premier temps retenir les modèles ayant une valeur de *box_loss* proche du palier inférieur

Puis, pour sélectionner un unique modèle, nous pouvons, parmi les modèles retenus, comparer le mAP50-95 (cf. Annexe VII), métrique synthétique des performances du modèle à l'issue de l'apprentissage. Selon ce critère, nous pouvons identifier l'un des modèles comme le plus performant, avec un mAP50-95 maximum à 0.24. Ainsi, le modèle minimisant le critère du *box_loss* (cf. Annexe VI) montre un mAP50-95 plus faible. Pour maximiser la métrique synthétique du mAP50-95, le modèle aux performances de *box_loss* très faiblement inférieur est sélectionné.

b. Performance du modèle sélectionné

Une fois le modèle cible identifié et entraîné avec les hyperparamètres testés aboutissant aux meilleures performances, on peut étudier la matrice de confusion des 3 classes d'objets à détecter obtenue en sortie de phase d'apprentissage (cf. Figure 10). La diagonale partant du coin en haut à gauche représente les vrais positifs ou détections pour lesquels le modèle a attribué la bonne classe et dont la localisation de la boîte englobante sur l'image est significativement semblable à la labellisation. Ainsi, 81 % des poux, 64 % des rougeurs et 44 % des pertes d'écaïlles sont bien détectés, selon les critères usuels d'évaluation des performances d'un modèle. Par ailleurs, on constate, pour certaines pathologies un grand nombre de faux négatifs, présents sur la ligne du bas : 56 % des pertes d'écaïlles et 36 % des rougeurs ne sont pas détectées par le modèle. Enfin, près de 21 % des prédictions pour les poux, 44 % pour les pertes d'écaïlles et 35 % pour les rougeurs sont des faux positifs. Par ailleurs, les confusions interclasses sont faibles ou nulles. Les erreurs de prédictions du modèle semblent porter davantage sur la détection et la position des boîtes englobantes que sur la classification des objets. C'est pourquoi nous n'avons pas développé davantage l'analyse des performances de classification du modèle, pour nous concentrer sur la détection.

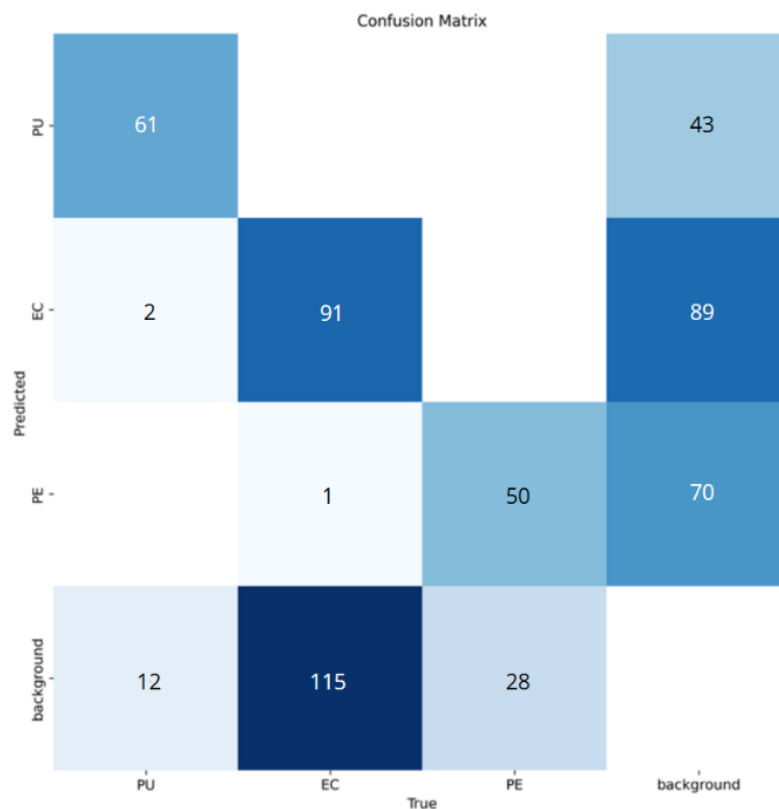


Figure 10 : Matrice de confusion des détections du modèle sélectionné (nombre d'objets par classe)

3. Evaluation des prédictions de la gravité

A partir des résultats obtenus par la méthode relative à la génération du code de gravité d'une pathologie, nous pouvons évaluer la capacité du modèle à prédire l'indice de gravité associé au code patho en fonction du seuil de confiance accordé aux détections pour chaque pathologie. Ces

performances suivent des tendances différentes en fonction du seuil. Pour toutes les pathologies, l'indice de gravité est de plus en plus sous-estimé en augmentant ce seuil (cf. Figure 11, 12 et 13). Par exemple, pour les pertes d'écaïlles (cf. Figure 11), le taux d'accord entre la prédiction et l'observation (delta de 0) est maximum au seuil de confiance à 0.01 (86 % de bonnes prédictions) et décroît si le seuil augmente. Pour les rougeurs (cf. Figure 12), on constate un maximum atteint au seuil de confiance de 0.01, signifiant que le modèle prédit à 94 % la bonne gamme de gravité de la pathologie, avec peu de sous-estimation. Quant aux poux, on observe une forte sous-estimation des indices de gravité prédits avec l'augmentation de la sélectivité des prédictions du modèle. Un plateau de meilleures performances est atteint à un seuil de confiance de 0.03 et 0.05 pour cette dernière pathologie (cf. Figure 13).

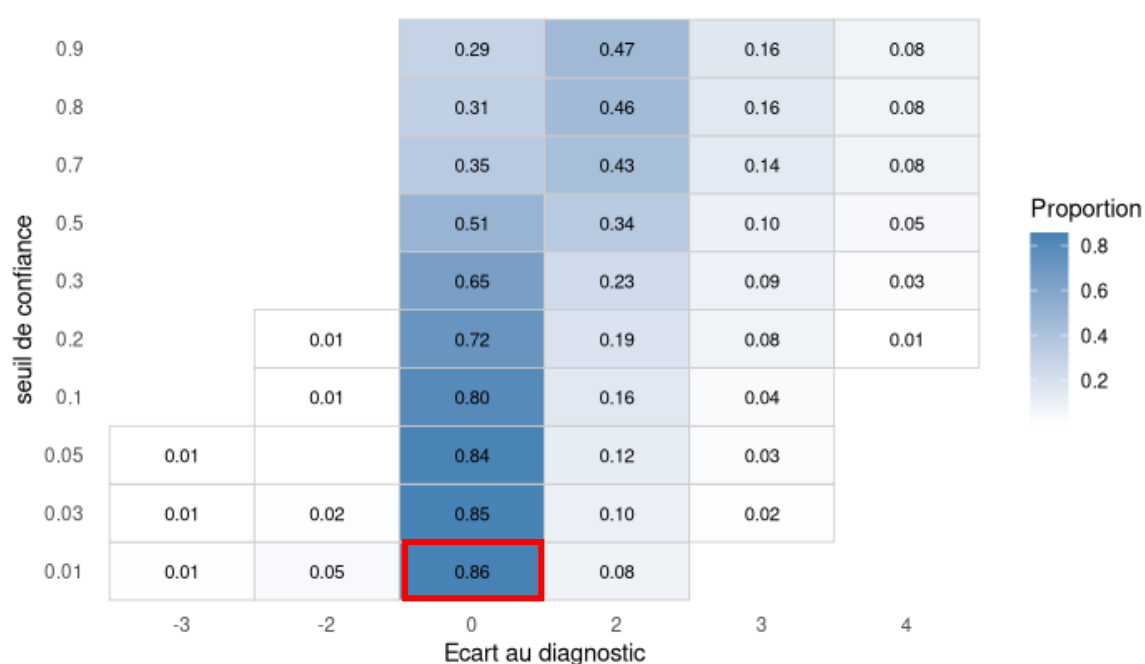


Figure 11 : Performances de prédiction de l'indice de gravité du modèle pour les pertes d'écaïlles

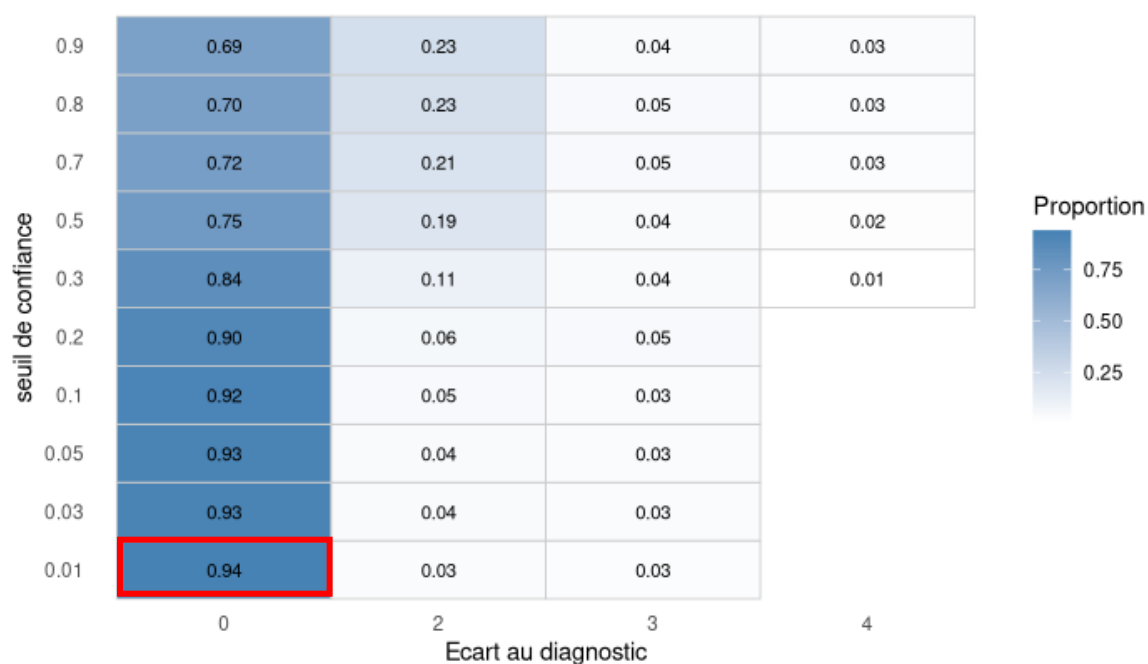


Figure 12 : Performances de prédiction de l'indice de gravité du modèle pour les rougeurs

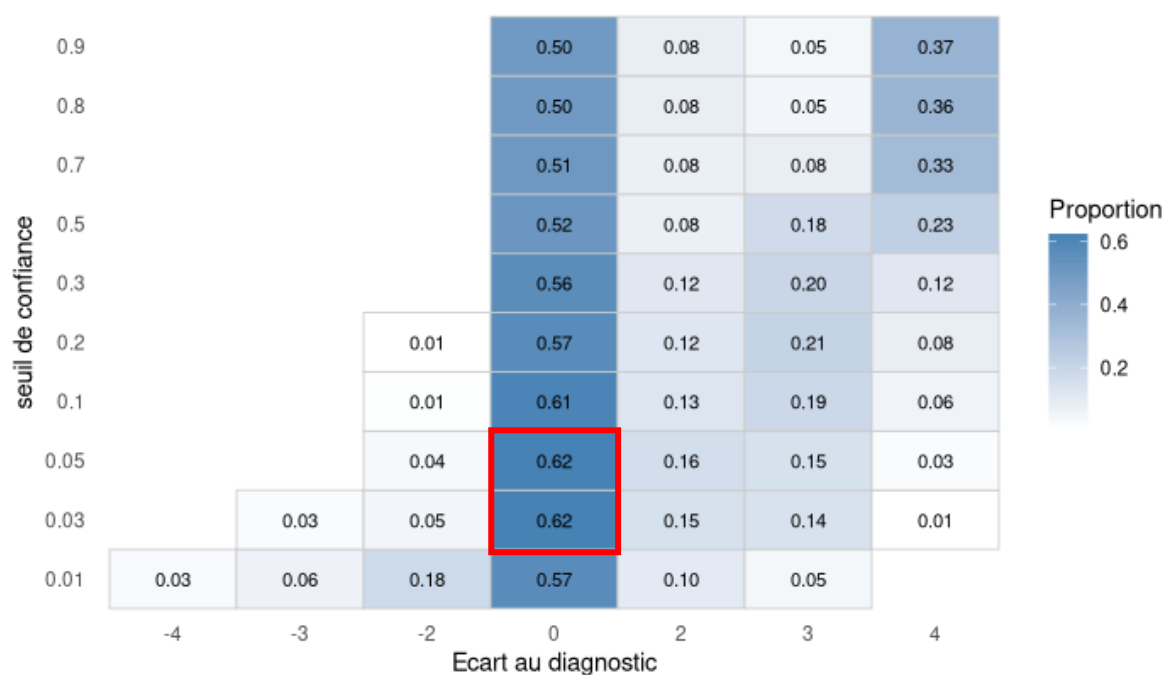


Figure 13 : Performances de prédiction de l'indice de gravité du modèle pour les poux

4. Evaluation des prédictions de la localisation

A partir des détections de chaque pathologie et leur segmentation sur l'image, nous avons obtenu un code de localisation associé à chaque pathologie. Selon le seuil de confiance accordé aux détections, nous pouvons étudier l'efficacité du pipeline à prédire le bon code de localisation. Le maximum est

atteint pour les seuils de confiance 0.03 pour les pertes d'écaïlles (avec une espérance d'utilité de 73 %), 0.2 pour les deux autres pathologies (avec 50 % pour les rougeurs et 60 % pour les poux) (cf. Figure 14).

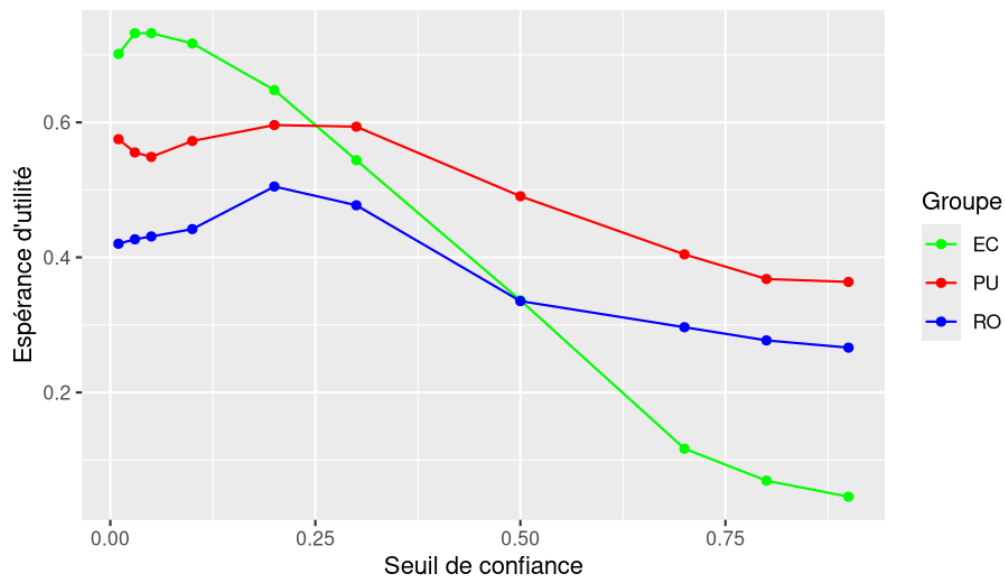


Figure 14 : Performances de prédiction du code de localisation du modèle pour les trois pathologies

5. Evaluation générale du modèle et le pipeline

Après avoir regardé les performances par tâche du pipeline (gravité, localisation), nous pouvons les comparer par pathologies, pour un même seuil de confiance pour identifier un seuil optimal. Ici, chaque point correspond à l'intersection entre les performances des deux tâches, par seuil de confiance. La ligne pointillée correspond à la fonction identité. Ainsi, si un point est supérieur à cette droite, le modèle est plus performant pour la prédiction de la gravité, pour un seuil et une pathologie donnée. De même, si le point est inférieur à cette droite, alors le modèle prédit mieux la localisation.

A partir de la figure 15, on peut identifier le seuil favorisant la qualité de prédiction de la localisation avec une performance très proche du maximum d'efficacité de prédiction de la gravité pour la classe des pertes d'écaïlles. En effet, au regard de la distance des points au point optimal de coordonnées (1, 1), le point du seuil à 0.03 est celui qui favorise les deux modalités au maximum. De même, le seuil de 0.2 maximise les performances du modèle pour les rougeurs (cf. Figure 16). Pour ces deux premières classes de pathologies, tous les points sont situés au-dessus de la droite : le modèle prédit, en moyenne, mieux la gravité que la localisation de ces pathologies. Pour la classe des poux (cf. Figure 17), certains seuils favorisent les performances portant sur la localisation (0.03 et 0.05) et d'autres sur la gravité (0.2). Les seuils qui maximisent les performances sont les valeurs de 0.03 et 0.2. Finalement, le seuil de confiance à 0.1 est le meilleur compromis entre les deux modalités de sélection du seuil de confiance, car il est le point le plus proche du point optimal de coordonnées (1, 1).

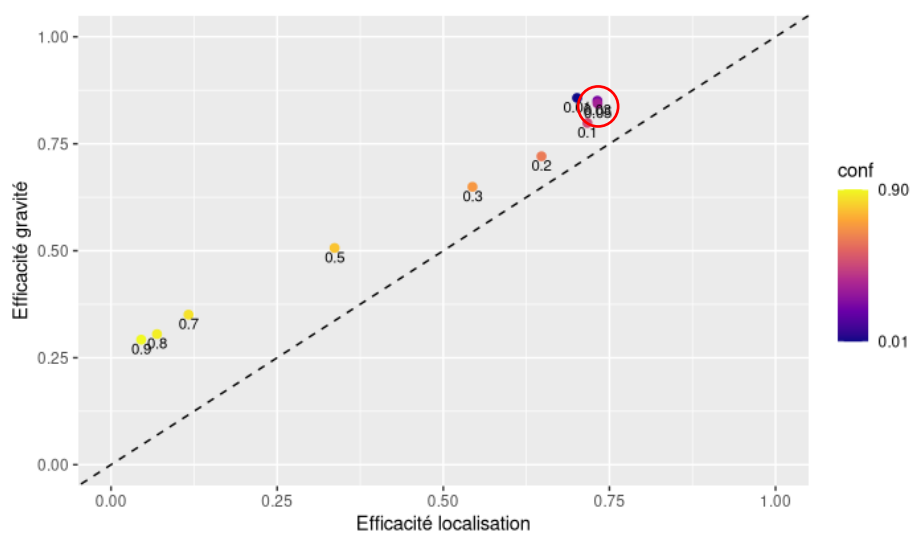


Figure 15 : Performances de prédiction de localisation et de gravité pour les pertes d'écailles

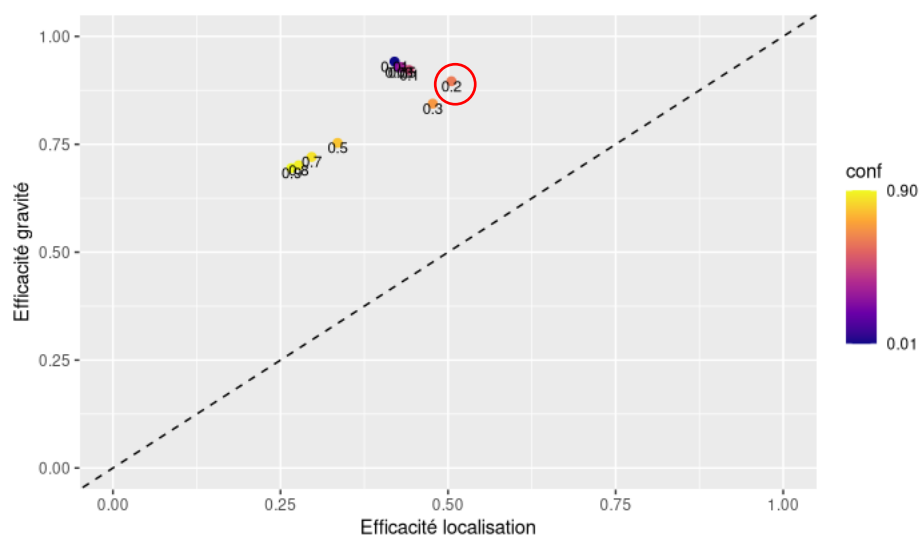


Figure 16 : Performances de prédiction de localisation et de gravité pour les rougeurs

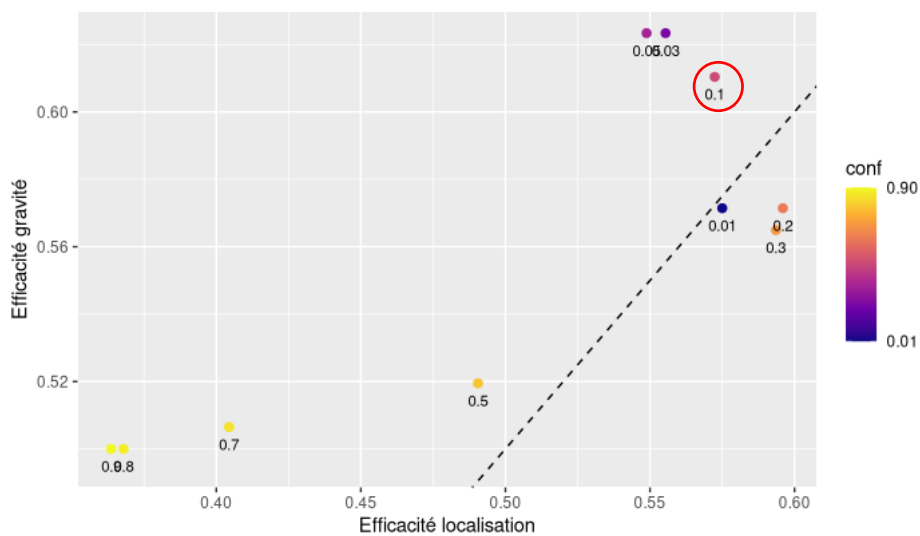


Figure 17 : Performances de prédiction de localisation et de gravité pour les poux

6. Précision, sensibilité et spécificité du modèle choisi

En considérant les valeurs de seuils de confiance optimaux ci-dessus, les performances du modèle de détection sont évaluées en termes de prédiction d'absence/présence de chaque pathologie. Les matrices de confusion, construites en comparant les prédictions aux observations précisent ces résultats.

Tableau 6 : Performances du modèle par pathologie aux seuils de confiance sélectionnés

	Seuil de confiance	Matrice de confusion	Précision	Sensibilité	Spécificité	Kappa
Pertes d'écaïlles	0.03	<pre> Réf. Positif Préd. Positif 7 143 Négatif 1 4 Réf. Négatif </pre>	0.93	0.97	0.13	0.12
Rougeurs	0.2	<pre> Réf. Positif Préd. Positif 15 80 Négatif 26 34 Réf. Négatif </pre>	0.68	0.70	0.63	0.29
Poux	0.1	<pre> Réf. Positif Préd. Positif 16 68 Négatif 40 31 Réf. Négatif </pre>	0.70	0.69	0.71	0.38

D'après les valeurs de sensibilité (cf. Tableau 6), le modèle, est performant pour détecter la présence d'objets sur les images. Pour autant, ce modèle est peu spécifique pour les pertes d'écaïlles (13 %), c'est-à-dire qu'il prédit avec beaucoup d'erreurs une absence effective de la pathologie, ce qui est moins le cas pour les rougeurs et poux pour lesquels le modèle prédit une bonne absence à 63 et 71 %. La précision du modèle pour chaque pathologie nous apprend qu'il est tout de même très performant en moyenne, sur le jeu de données test pour les pertes d'écaïlles, avec un taux global de bonnes réponses de 93 %. Ses performances moyennes diminuent à 68 et 70 % pour les deux autres

pathologies (rougeurs et poux, respectivement). Par ailleurs, le Kappa de Cohen, compris entre 0.12 et 0.38, montre que le modèle prédit mieux qu'un modèle aléatoire. Ce coefficient valide qu'au moins une partie des caractéristiques visuelles des pathologies sont automatiquement identifiées par le modèle, lui permettant de distinguer les individus sains et malades pour chaque pathologie. Le Kappa tient compte du déséquilibre entre les observations de positifs et négatifs. Dans notre jeu de données test, la majorité des poissons est atteinte par une ou plusieurs pathologies, ce qui biaise l'interprétation de la précision, accord moyen global. C'est pourquoi le Kappa de Cohen est le plus bas pour les pertes d'écaillés (0.12), alors que la précision est maximale (93 %).

7. Exemple d'application

L'intégralité du pipeline est déployée sur les photos prises au cours de l'année 2025 (sur 125 poissons), pour estimer le nombre de poissons infectés par les poux de mer. Les résultats du modèle sont comparés aux évaluations des opérateurs, notées lors de la capture des individus. Les poux ayant tendance à se décrocher à l'arrivée des salmonidés en eau douce, leur présence peut être considérée comme un indicateur du temps de séjour en eau douce des poissons migrateurs. Comme évoqué plus haut, le seuil de confiance de sélection est fixé pour la classe des poux à 0.1, signifiant que toutes les détections ayant un score de confiance supérieur sont conservées (cf. Figure 18).



Figure 18 : Exemple de détection des poux sur le corps d'un poisson

En moyenne, sur la saison, 67 % des individus sont atteints par les poux de mer et notre modèle prédit une proportion de 71 % (cf. Figure 19). On constate que sur les semaines regroupant le plus d'effectif (semaines 21, 23 et 26), les taux de poissons infectés prédits sont très proches des taux observés. Cependant, sur les plus petits effectifs, la différence est plus importante.

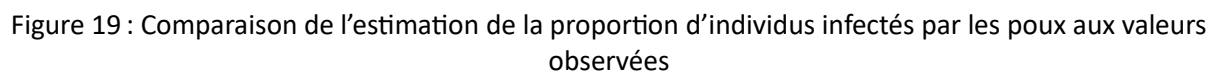


Figure 1 consists of two stacked bar charts, (a) and (b), showing the number of observations (Nombre d'observations) on the y-axis (ranging from 0 to 25) against the week (Semaine) on the x-axis (May, June, July). The legend for both charts indicates three types of events (Type) and three levels of gravity (Gravité).

Chart (a): Shows the number of observations for Type 0 (white), Type 1-2 (light blue), and Type 3-4 (dark blue). The data shows a significant increase in observations in late June and early July, with Type 3-4 being the most frequent category in those weeks.

Chart (b): Shows the number of observations for Type 0 (white), Type 1-2 (light blue), and Type 3-4 (dark blue). The data shows a significant increase in observations in late June and early July, with Type 1-2 being the most frequent category in those weeks.

Figure 20 : Comparaison des histogrammes du nombre d'individus atteints par les poux par semaine
(a) observés, (b) prédits

Discussion

1. Phase de labellisation

L'un des objectifs de l'utilisation de modèles d'analyse d'images par apprentissage profond est d'automatiser la tâche de détection et classification des pathologies. L'implication de certains biais cognitifs et d'identification de la pathologie semble inévitable dans une méthode impliquant l'intervention d'un opérateur humain (Taylor et al. 2011). Pour autant, le modèle de détection est construit par un apprentissage supervisé, dont les données d'entrée ont été labellisées. Il se base donc sur des données, préalablement renseignées par une personne, et cette phase est sujette à la subjectivité opérateur-dépendante (Rädsch et al. 2023). Ainsi, la position d'une boîte et parfois même l'appartenance d'une pathologie à une classe labellisée peut dépendre du labellisateur. Au cours de la phase de labellisation par les experts, certaines sources de discordances ont pu être identifiées et rectifiées selon les règles suivantes : ne pas labelliser une pathologie qui n'aurait pas été prise en compte dans l'analyse sur le terrain, ou encore englober toute la pathologie, bien que diffuse, tant qu'il y a une continuité de celle-ci sur le corps du poisson. Ce dernier point permet au modèle d'étudier des grosses unités d'objets et non des petits patches discontinus.

Ces discordances de choix de labellisation ou non d'une pathologie selon l'observateur est difficilement contournable. La méthode de pointage de la pathologie, quant à elle, pourrait être affinée par une phase de labellisation impliquant d'autres informations. Nous pourrions considérer une labellisation sous forme de masque et non de boîtes englobantes, ou une méthode hybride compilant les boîtes et un masque de contour des objets, affinant ainsi la surface concernée et améliorant les performances du modèle (Kang, Kim 2019). Cependant, cette méthode nécessite une attention et un temps de traitement plus long, charge de travail que les experts sur les stations de contrôles ne pouvaient assumer à cette saison.

De plus, l'amélioration des performances du modèle de détection des pathologies aurait pu être appuyé par des méthodes de *transfer-learning* depuis des modèles plus spécifiques. Cependant, l'accès aux ressources de modèles réalisant des tâches similaires, parfois appliqués à un contexte d'aquaculture, est très limité (Ahmed, Aurpa, Azad 2022; Zhang et al. 2024).

2. Partitionnement du poisson

L'identification des zones du corps du poisson concernées par chaque pathologie est très dépendante de la méthode appliquée pour identifier chaque point d'intérêt. Entre autres, le ratio de la longueur totale pour positionner l'anus n'est pas identique selon l'espèce. Dans notre jeu de données, nous n'avons considéré que deux espèces de salmonidés morphologiquement proches, mais ce ratio diffère légèrement. En effet, la longueur pré-anale est d'environ 69 % de la longueur totale pour *Salmo salar* et de 66 % pour *Salmo trutta* (FishBase 2025b; 2025a). L'intégration de l'espèce dans ce pipeline permettrait de tenir compte de spécificités intrinsèques à chaque espèce pour affiner nos résultats et ainsi étendre notre méthode à d'autres espèces, morphologiquement différentes, ou encore à des stades de développement spécifiques. Pour autant, aujourd'hui ce découpage

morphologique est réalisé à vue et l'attribution du code de localisation ne semble pas, d'après les opérateurs, à 3 % près de la longueur du poisson.

De plus, le résultat du partitionnement morphologique est dépendant de la méthode d'obtention du sens et de l'orientation de l'individu sur l'image. Dès qu'un objet n'est pas, ou est mal détecté, alors le code patho de localisation sera erroné et probablement invalidé. Une autre méthode pour obtenir les points de la bouche et la fourche, aurait été d'entraîner un modèle de détection de points comme Yolo11.pose (Tseng, Hsieh, Kuo 2020). Cet outil n'a pas été retenu car aurait nécessité une autre phase de labellisation chronophage, tandis que le pipeline d'extraction des points d'extrémités a été mis en place lors d'une phase préliminaire du stage, ayant pour objectif d'appréhender les outils d'analyse d'images.

En plus de certaines images non exploitables par les modèles utilisés dans le partitionnement du poisson car ils manquent un objet, une autre partie des images est omise à cause de leur qualité qui ne permet pas de les analyser. Par exemple, lorsqu'une main de l'opérateur recouvre une partie du poisson et les prises de vues détails d'une partie du corps, ne présentant pas l'ensemble du poisson sur le flan comme le dos, l'anus ou la tête. Comme des pathologies ne sont pas analysées par la méthode, certains codes pathologie sont sous-estimés et demanderaient une étude plus exhaustive du corps.

3. Evaluation du modèle de détection

Notre évaluation du pipeline, au regard de la prédiction de la localisation et de la gravité, ainsi que de la prédiction de l'absence/présence de chaque pathologie, semble indiquer que la méthode est fonctionnelle. En effet, la simple prise en compte des métriques d'évaluations de performances usuelles pour un modèle de détection n'est pas pertinente pour comprendre les performances opérationnelles du modèle Yolo. Cela pourrait s'expliquer par une invalidation trop stricte des détections prédites, qui, *a posteriori*, permettront tout de même d'identifier des pathologies.

En outre, chaque pathologie peut être associée à une ou plusieurs problématiques spécifiques. Par exemple, les poux sont de très petits objets dont la forme n'est pas strictement identique selon son stade de développement. A faible résolution, le modèle ne peut extraire que peu d'informations de ces petits objets (Lin et al. 2017). Tandis que les rougeurs et pertes d'écailles sont des tâches parfois diffuses ou peu contrastées sur le corps du poisson. Il n'est donc pas aisé d'extraire des entités cohérentes pour le modèle lors de sa détection (Cho et al. 2022; Adhikari et al. 2024). Pour ces objets, les contours ne sont pas toujours clairement définis et sont morcelés par patch, ce qui rend la détection et la validation complexe. Il faut donc baser l'entraînement du modèle de détection sur un très grand jeu de données pour améliorer sa généralisation, avant d'atteindre un plateau de performances (Gottlich et al. 2023).

4. Evaluation de la prédiction du code patho

L'évaluation du pipeline à prédire les codes patho a montré qu'à partir du seuil de confiance à 0.01, on observe une monotonie de l'efficacité de prédiction de l'indice de gravité pour les pertes d'écailles et les rougeurs. Nous aurions pu nous attendre à une croissance de ces performances jusqu'à un certain seuil, du fait de la sélection de détections plus pertinentes, en deçà duquel, nous surestimons la gravité

en conservant trop de boîtes englobantes. Cependant, pour observer cette tendance, il nous faudrait abaisser le seuil de confiance minimum de notre analyse, déjà très bas ici. Ce manque de sélectivité des détections nécessite de valider, par pathologie, les détections par le code de localisation. Cela permet de confirmer la quantité de détection retenue et leur place sur l'image, alternative aux métriques usuelles d'évaluation des performances.

Pour le cas des poux, on constate une diminution de la surestimation et une augmentation de sous-estimations avec l'accroissement du seuil de confiance. Cela s'explique par une restriction supérieure quant à la sélection des détections conservées pour l'estimation de la gravité, jusqu'à omettre un trop grand nombre de boîtes et ainsi sous-estimer la gravité de la pathologie. Notre choix du seuil de confiance pourra donc se tourner vers le score maximum de prédictions justes, privilégier un modèle exhaustif au risque de surestimer la gravité ou adopter un seuil plus important pour être davantage précautionneux mais au risque de sous-estimer la gravité.

Ces tendances se retrouvent dans l'évaluation de la prédiction de la localisation. Plus le seuil augmente, plus on affine la sélection des boîtes dans un premier temps et donc cela rapproche le résultat des observations. Par la suite, les détections sont de moins en moins nombreuses et omettent une partie du recouvrement de la pathologie sur le corps du poisson. Par ailleurs, l'efficacité de la localisation des pertes d'écaïlles est inférieure à celles des deux autres classes. Cela peut s'expliquer par un nombre très limité d'individus non concernés par les pertes d'écaïlles. Il y a donc un grand nombre de mauvaises localisations, notamment en cas d'absence pour les pertes d'écaïlles, compensé par des absences bien détectées à des hauts seuils de confiances pour les autres pathologies. De plus, si, dans le cas de la note moyenne de localisation pour les poux, les performances ont une allure décroissante, puis croissante et de nouveau décroissante, cela peut être expliqué par le système de notation des prédictions de localisation. En effet, nous avons fait le choix de hiérarchiser les erreurs avec une plus forte pénalisation des prédictions trop larges, par rapport aux prédictions incluses dans le code observé. Il faut donc bien exploiter la note de localisation par pathologie et la comparer relativement entre seuils de confiance sans accorder d'interprétation quantitative. La localisation ne semble pas au cœur de l'attention portée aux pathologies, mais permet de confirmer ou non un code patho afin d'en étudier la gravité.

5. Choix du seuil de confiance

L'analyse des résultats conjoints des performances de prédiction de la gravité et de la localisation nous permet d'avoir une vision synthétique des deux composantes prédites par le pipeline afin de sélectionner le meilleur seuil de confiance des détections du modèle Yolo, par pathologie. La gravité est une performance quantitative tandis que la localisation permet de valider ou non les détections. Il est donc indispensable de considérer un compromis entre gravité et localisation, en se rapprochant au maximum du point sur les graphiques figures 15, 16 et 17 de coordonnées (1, 1). Dans le cas du choix du seuil de confiance pour les rougeurs, le seuil à 0.01, maximise les performances de prédiction de la gravité. Cependant, moyennant une baisse de 4 % de bonnes prédictions de la gravité, nous augmentons de 9 % l'efficacité de localisation, validant ainsi davantage les positions des boîtes détectées. Il en va de même pour le choix du seuil de confiance à 0.1 plutôt que 0.03 ou 0.2 pour les poux. A ces valeurs de seuil de sélection des détections par leur score de confiance, on ne peut pas

considérer que le modèle soit confiant dans ces prédictions mais suffisamment pour les valider et les étudier afin d'obtenir le code patho.

On constate également que le choix du seuil de confiance pour les pertes d'écaïlles et les rougeurs maximisent systématiquement la gravité par rapport à la localisation. On pourrait expliquer cela par la nature diffuse de ces pathologies, qui permettent difficilement de statuer sur leur localisation exacte. D'un autre côté, les poux, petits objets sur l'image, favorisent la prédiction de la gravité en manquant quelques individus, mais pouvant fortement impacter le code de localisation, si les poux manqués sont sur une autre partie du corps du poisson.

6. Précision, sensibilité et spécificité du modèle de prédiction de présence/absence

Une fois le seuil de confiance fixé, nous pouvons affiner l'évaluation du modèle Yolo de détection des pathologies en étudiant sa capacité à prédire les présences et absences de chaque pathologie sur les individus.

On remarque, tableau 6, que les matrices de confusion sont déséquilibrées, avec une représentativité très forte de cas positifs par rapport aux négatifs. Cela est dû aux codes pathos renseignés dans le jeu de données disponible (cf. Figure 5). Que ce soit pour le jeu de données d'entraînement ou pour l'évaluation, peu d'individus sont renseignés comme sains dans la base.

Avec une précision comprise entre 0.68 et 0.93, nous pouvons considérer que le modèle a une capacité très différente de prédire le bon état du poisson (sain ou affecté) en fonction des pathologies étudiées. Le modèle atteint de très bonnes performances pour les pertes d'écaïlles, principalement influencée par la forte représentativité des cas positifs, et celles-ci sont correctes pour les deux autres pathologies. Ce déséquilibre se constate également au travers d'un très bon score de sensibilité pour les pertes d'écaïlles par rapport à la spécificité. Ces résultats soulignent la nécessité d'ajouter des images de poissons sains au jeu d'entraînement et d'évaluation si l'on souhaite rééquilibrer la capacité du modèle à être précautionneux et exhaustif. Ce qui est davantage le cas pour les deux autres pathologies.

De plus, les coefficients du Kappa de Cohen sont inférieurs à 0.5, indiquant que l'accord entre les prédictions et observations est faible. Pour autant, le modèle semble mieux faire que le hasard ($Kappa < 0$) indiquant qu'il a tout de même extrait des caractéristiques intéressantes de chaque pathologie. La principale cause de ce Kappa faible vient du déséquilibre entre les individus infectés ou non et les capacités limitées du modèle en la prédiction de cas négatifs (Johnson, Khoshgoftaar 2019).

7. Application à un cas concret

Les résultats sur la détection de l'absence/présence des poux sur les individus de l'année 2025 semblent correspondre raisonnablement aux observations, conformément à l'analyse de sensibilité et spécificité menée plus haut. Cependant, le modèle sous-estime quasi systématiquement la gravité, lorsque la pathologie est présente, signifiant qu'il ne comptabilise pas tous les poux sur le poisson. En effet, seules les images des flancs de l'individu sont analysées car les poux ne sont pas détectés sur le

dos du poisson, par manque de contraste. Il pourrait donc être intéressant d'isoler les images du détail du dos comportant des poux afin d'entraîner un modèle spécifique et affiner notre résultat. La caractérisation des images du dos et de l'anus, sur lesquelles on retrouve fréquemment des poux, pourraient être prises en compte en les compilant dans le code patho du haut du corps (HT) pour le dos et bas du corps pour l'anus (AT). Ainsi, nous apporterions une correction de la prédiction du code patho, vers une augmentation de l'estimation de la gravité.

Conclusion

Au cours de ce stage, nous avons développé une méthode automatique d'obtention des codes pathologies, utilisés par les opérateurs des stations de contrôle des poissons migrateurs pour décrire finement les pathologies présentes sur le corps des individus. L'utilisation de différents algorithmes d'analyse d'images par apprentissage profond nous a permis d'identifier les pathologies sur le poisson et en extraire les informations de gravité et localisation, nécessaires à la construction des codes patho. Ces outils d'intelligence artificielle, couplées à des stratégies en post-traitement, sont intégrées au pipeline, réalisant de nombreuses tâches, telles que la détection des pathologies, l'estimation de la courbure et de la longueur fourche du poisson, son découpage morphologique et la segmentation des pathologies pour connaître leur surface d'expansion. Dans le cadre de cette étude, nous avons tiré bénéfice des spécificités de chaque modèle utilisé tel qu'un apprentissage efficace et rapide à prendre en main pour Yolo, une capacité de détection de nombreux objets sans apprentissage sur un jeu de données spécifique à notre étude pour Grounding Dino, ou encore une segmentation d'objets, encore une fois sans apprentissage de notre part pour SAM.

Tout ceci nous a amené à prédire des codes patho dont l'évaluation a montré des résultats différents selon les pathologies, expliqués principalement par la nature visuelle des pathologies et la composition de notre jeu de données. En effet, les prédictions des pertes d'écailles présentent de très bons résultats lorsque la pathologie est présente, avec moins de 3 % de faux négatifs et 85 % des prédictions de gravité sont considérées comme acceptables. Les performances du pipeline pour les deux autres pathologies sont plus équilibrées pour prédire les absences et présences et 90 % de bonnes prédictions de la gravité pour les rougeurs et 57 % pour les poux de mer.

En effet, le modèle atteint des performances de prédictions des codes pathologie intéressantes, la chaîne de traitement nécessiterait d'être complétée par des modèles spécifiques de détections des unités morphologiques des poissons ainsi que par un grand nombre de nouvelles données photographiques. De plus, comme évoqué dans la présentation des modèles Yolo, les modèles d'analyse d'images sont de plus en plus efficaces et les puissances de calculs des outils numériques sont également croissantes, ce qui nous amène à penser que ce type de méthodes pourra voir son efficacité et son accessibilité améliorées.

Perspectives

Dans le cadre de ce projet, nous pouvons projeter d'avoir accès à davantage d'images dans le futur, issues de la Bresle, voire d'autres rivières contrôlées. Cet enrichissement de la base de données permettrait de palier à plusieurs problèmes tel que le déséquilibre entre les individus sains et affectés par chaque pathologie. Nous pouvons également espérer une augmentation du nombre d'images et de leur variabilité à analyser par le modèle lors de son apprentissage, allant dans le sens d'une généralisation des performances de détection des différentes classes. La couverture de l'ensemble des saisons ouvrirait la possibilité de distinguer différents modèles de détection selon la saison, lorsque la robe des individus diffère largement. Ces images pourront également alimenter d'autres modèles d'identifications de la nageoire caudale, de la tête ou autres points d'intérêts comme le point fourche, l'extrémité de la bouche, améliorant la précision des différentes tâches du pipeline d'obtention des codes patho.

Le protocole d'acquisition des données photographiques et descriptives des pathologies, mis en place avec succès sur la Bresle, pourrait être étendu aux différentes stations de contrôles des poissons migrateurs amphihalins de la façade atlantique (Scorff, Nivelle) ou de la Manche (Oir) inclus dans l'ORE DiaPFC, afin d'acquérir de nouvelles images. De plus la présentation aux opérateurs des résultats préliminaires a éveillé leur curiosité et répondu à leurs questions sur les applications possibles des outils d'intelligences artificielles, intégrés à leurs activités. Cela a notamment alimenté des discussions sur l'installation d'une caméra sur une potence pour faciliter et uniformiser la prise de photographies des salmonidés lors de leur biométrie. Ce dispositif offrirait une meilleure standardisation des images grâce à un focus, un cadrage et une luminosité adaptée.

A l'issue de cette phase exploratoire, nous pouvons envisager, en accord avec les opérateurs des stations de contrôle, d'intégrer ces outils de génération des codes patho à leur protocole d'acquisition des données biométriques. Une phase de mise en forme d'une interface opérationnelle, intégrée à la table de biométrie permettrait aux agents de prendre en photo chaque individu et de l'analyser par la méthode développée au cours de notre étude. Ainsi, des codes patho pourraient être générés et supervisés par les experts en temps réel, comme c'est déjà le cas pour la taille ou la masse des poissons. Un tel outil d'automatisation de l'acquisition des informations pathologiques pourrait permettre de réduire le temps de manipulation des poissons hors de l'eau, objectif clairement identifié dans le cadre du Bien-Etre Animal et de s'affranchir de certains biais humains liés à l'estimation des indices de gravités (particulièrement pour le recouvrement des pathologies diffuses).

Bibliographie

ADHIKARI, S., KUMAR, R., MOPURI, K. R. et PACHAMUTHU, R., 2024. Lost in context: The influence of context on feature attribution methods for object recognition. In : *Proceedings of the Fifteenth Indian Conference on Computer Vision Graphics and Image Processing* [en ligne]. 13 décembre 2024. pp. 1-10. [Consulté le 5 août 2025]. Disponible à l'adresse : <http://arxiv.org/abs/2411.02833>

AHMED, Md. S., AURPA, T. T. et AZAD, Md. A. K., 2022. Fish disease detection using image based machine learning technique in aquaculture. *Journal of King Saud University - Computer and Information Sciences*. septembre 2022. Vol. 34, n° 8, pp. 5170-5182. DOI 10.1016/j.jksuci.2021.05.003.

ARKIN, E., YADIKAR, N., XU, X., AYSA, A. et UBUL, K., 2023. A survey: object detection methods from CNN to transformer. *Multimedia Tools and Applications*. juin 2023. Vol. 82, n° 14, pp. 21353-21383. DOI 10.1007/s11042-022-13801-3.

BAKKE, T. A. et HARRIS, P. D., 1998. Diseases and parasites in wild Atlantic salmon (*Salmo salar*) populations. *Canadian Journal of Fisheries and Aquatic Sciences*. 1998. Vol. 55, n° S1. DOI 10.1139/d98-021.

BELPAIRE, C. et GOEMANS, G., 2007. The european eel *Anguilla anguilla*, a rapporteur of the chemical status for the water framework directive? *Vie et Milieu*. 2007. pp. 235-252.

BERTRAM, C. A. et KLOPFLEISCH, R., 2017. The Pathologist 2.0: An update on digital pathology in veterinary medicine. *Veterinary Pathology*. septembre 2017. Vol. 54, n° 5, pp. 756-766. DOI 10.1177/0300985817709888.

BLONDIAUX, E., SILEO, C., DOHNA, M. et CHALARD, .F, 2018. *Etiopathogénie des erreurs (et comment y remédier...)* [en ligne]. 2018. Disponible à l'adresse : https://www.sfip-radiopédiatrie.org/wp-content/uploads/2018/07/blondiaux_trousseau_2017.pdf

BOLSTAD, G. H., DISERUD, O. H., PATERSON, R. A., ULVAN, E. M., KARLSSON, S., UGEDAL, O. et NÆSJE, T. F., 2025. A fitness-based indicator for the effect of aquaculture-produced salmon lice on wild sea trout. BYRON, C. (éd.), *ICES Journal of Marine Science*. 1 mai 2025. Vol. 82, n° 5, pp. fsae192. DOI 10.1093/icesjms/fsae192.

BUCKE, D., VETHAAK, D., LANG, T. et S. MELLERGAARD, S., 1996. 19 : *Common diseases and parasites of fish in the North Atlantic: Training guide for identification* [en ligne]. Training Guide. ICES. Disponible à l'adresse : https://www.researchgate.net/publication/287816176_Common_diseases_and_parasites_of_fish_in_the_North_Atlantic_Training_guide_for_identification?enrichId=rgreq-da5d23dd6657f3f7246eff594dbe2295-XXX&enrichSource=Y292ZXJQYWdlOzI4NzgxNjE3NjBUzozMDk1NjY4MTI2ODgzODdAMTQ1MDgxNzgz4NzcxNQ%3D%3D&el=1_x_2&esc=publicationCoverPdf

CHO, S.-J., KIM, S.-W., JUNG, S.-W. et KO, S.-J., 2022. Blur-robust object detection using feature-level deblurring via self-guided knowledge distillation. *IEEE Access*. 2022. Vol. 10, pp. 79491-79501. DOI 10.1109/ACCESS.2022.3194898.

COSTELLO, M. J., 2009. How sea lice from salmon farms may cause wild salmonid declines in Europe and North America and be a threat to fishes elsewhere. *Proceedings of the Royal Society*

B: *Biological Sciences*. 7 octobre 2009. Vol. 276, n° 1672, pp. 3385-3394. DOI 10.1098/rspb.2009.0771.

DAVELAAR, E. J., TIAN, X., WEIDEMANN, C. T. et HUBER, D. E., 2011. A habituation account of change detection in same/different judgments. *Cognitive, Affective, & Behavioral Neuroscience*. décembre 2011. Vol. 11, n° 4, pp. 608-626. DOI 10.3758/s13415-011-0056-8.

DEVLIN, J., CHANG, M.-W., LEE, K. et TOUTANOVA, K., 2019. BERT: Pre-training of deep bidirectional transformers for language understanding. . 2019. DOI 10.48550/arXiv.1810.04805.

DIRM - MANCHE EST-MER DU NORD, 2025. Arrêté - *DIRM - Interdiction temporaire de la pêche marine des saumons (Salmo salar) dans les eaux maritimes de la région Normandie*. 20 janvier 2025.

DUBOIS, R., 2024. *Taxonomie et morphome frie, peut-on remplacer l'humain par la machine ?* Rapport de stage (M2). Master Modélisation en écologie.

EBBINGHAUS, H., 2013. Memory: A contribution to experimental psychology. RUGER, H. A. et BUSSENIUS, C. E. (trad.), *Annals of Neurosciences*. 2013. Vol. 20, n° 4, pp. 155-156. DOI 10.5214/ans.0972.7531.200408.

EGEVAD, L., CHEVILLE, J., EVANS, A. J., HÖRNBLAD, J., KENCH, J. G, KRISTIANSEN, G., LEITE, K. R. M., MAGI-GALLUZZI, C., PAN, C.-C., SAMARATUNGA, H., SRIGLEY, J. R., TRUE, L., ZHOU, M., CLEMENTS, M., DELAHUNT, B., et THE ISUP PATHOLOGY IMAGEBASE EXPERT PANEL, 2017. Pathology Imagebase - A reference image database for standardization of pathology. *Histopathology*. novembre 2017. Vol. 71, n° 5, pp. 677-685. DOI 10.1111/his.13313.

ELIE, P. et GIRARD, P., 2007. *Manuel d'identification des principales lésions anatomo-morphologiques et des principaux parasites externes des anguilles européennes (Anguilla anguilla)* [en ligne]. 2007. Disponible à l'adresse : https://www.researchgate.net/publication/281443647_Manuel_d%27identification_des_principales_lesions_anatomo-morphologiques_et_des_principaux_parasites_externes_des_anguilles_europeennes_Anguilla_anguilla?enrichId=rgreq-92a93bd23248ce4208ce52a29a29641e-XXX&enrichSource=Y292ZXJQYWdIOzI4MTQ0MzY0NztBUzoyNjkzMzAwNTUxNjgwMDBAMTQ0MTlyNDY5NzUyNg%3D%3D&el=1_x_2&esc=publicationCoverPdf

ELIE, P. et GIRARD, P., 2014. *La santé des poissons sauvages : les codes pathologie, un outil d'évaluation*. Association Sané Poissons Sauvages. Association Santé Poissons Sauvages.

ENESCU, R., 2009. Effets sériels et conflit cognitif dans l'administration des moyens de preuves et le choix d'un verdict pénal. [en ligne]. 2009. Disponible à l'adresse : https://www.academia.edu/976185/Effets_S%C3%A9riels_et_Conflit_Cognitif_dans_L_Administration_des_Moyens_de_Preuves_et_le_Choix_d_un_Verdict_P%C3%A9nal

FAUSCH, K. D., KARR, J. R. et YANT, P. R., 1984. Regional application of an index of biotic integrity based on stream fish communities. *Transactions of the American Fisheries Society*. janvier 1984. Vol. 113, n° 1, pp. 39-55. DOI 10.1577/1548-8659(1984)113%3C39:RAOAI0%3E2.0.CO;2.

FISHBASE, 2025a. FishBase, *Salmo trutta*. *FishBase* [en ligne]. 2025. [Consulté le 23 juin 2025]. Disponible à l'adresse : <https://www.fishbase.us/summary/Salmo-trutta.html>

FISHBASE, 2025b. FishBase, *Salmo salar*. *FishBase* [en ligne]. 2025. [Consulté le 12 mai 2025]. Disponible à l'adresse : <https://www.fishbase.us/summary/Salmo-salar.html>

GOTTLICH, H. C., GREGORY, A. V., SHARMA, V., KHANNA, A., MOUSTAFA, A. U., LOHSE, C. M., POTRETZKE, T. A., KORFIATIS, P., POTRETZKE, A. M., DENIC, A., RULE, A. D., TAKAHASHI, N., ERICKSON, B. J., LEBOVICH, B. C. et KLINE, T. L., 2023. Effect of dataset size and medical image modality on convolutional neural network model performance for automated segmentation: A CT and MR renal tumor imaging study. *Journal of Digital Imaging*. 17 mars 2023. Vol. 36, n° 4, pp. 1770-1781. DOI 10.1007/s10278-023-00804-1.

GROUNDSED-SAM CONTRIBUTORS, 2023. *Grounded-Segment-Anything* [en ligne]. Jupyter Notebook. avril 2023. [Consulté le 14 avril 2025]. Disponible à l'adresse : <https://github.com/IDEA-Research/Grounded-Segment-Anything>

HÅRSTEIN, T et BULLOCK, G., 1976. An acute septicaemic disease of brown trout (*Salmo trutta*) and Atlantic salmon (*Salmo salar*) caused by a *Pasteurella* -like organism. *Journal of Fish Biology*. janvier 1976. Vol. 8, n° 1, pp. 23-26. DOI 10.1111/j.1095-8649.1976.tb03903.x.

HEFFERNAN, M. L., 1999. A review of the ecological implications of mariculture and intertidal harvesting in Ireland. *Irish Wildlife Manuals* [en ligne]. 1999. N° 7. Disponible à l'adresse : <http://hdl.handle.net/2262/69387>

ICES, 2024a. *Atlantic salmon (Salmo salar) from the Northeast Atlantic* [en ligne]. ICES Advice: Recurrent Advice. [Consulté le 7 juillet 2025]. Disponible à l'adresse : https://ices-library.figshare.com/articles/report/Atlantic_salmon_i_Salmo_salar_i_from_the_Northeast_Atlantic/25019639

ICES, 2024b. *Atlantic salmon (Salmo salar) from the Northeast Atlantic* [en ligne]. ICES Advice: Recurrent Advice. [Consulté le 14 avril 2025]. Disponible à l'adresse : https://ices-library.figshare.com/articles/report/Atlantic_salmon_i_Salmo_salar_i_from_the_Northeast_Atlantic/25019639

JENSEN, K.W., 1968. Sea trout (*Salmo trutta* L.) of the River Istra, western Norway. In : *Report of the Institute of Freshwater Research, Drottningholm* [en ligne]. Institute of Freshwater Research. pp. 187-213. Disponible à l'adresse : <https://gupea.ub.gu.se/handle/2077/48610>

JOCHER, G., QIU, J. et CHAURASIA, A., 2023. *Ultralytics YOLO* [en ligne]. Python. janvier 2023. [Consulté le 14 avril 2025]. Disponible à l'adresse : <https://github.com/ultralytics/ultralytics>

JOHNSON, J. M. et KHOSHGOFTAAAR, T. M., 2019. Survey on deep learning with class imbalance. *Journal of Big Data*. décembre 2019. Vol. 6, n° 1, pp. 27. DOI 10.1186/s40537-019-0192-5.

JOYCE, Edward J. et BIDDLE, Gary C., 1981. Anchoring and Adjustment in Probabilistic Inference in Auditing. *Journal of Accounting Research*. 1981. Vol. 19, n° 1, pp. 120. DOI 10.2307/2490965.

KANG, B. R. et KIM, H. Y., 2019. *BshapeNet: Object detection and instance segmentation with bounding shape masks* [en ligne]. 31 juillet 2019. arXiv. arXiv:1810.10327. [Consulté le 4 août 2025]. Disponible à l'adresse : <http://arxiv.org/abs/1810.10327>

KARR, J. R., 1981. Assessment of biotic integrity using fish communities. *Fisheries*. novembre 1981. Vol. 6, n° 6, pp. 21-27. DOI 10.1577/1548-8446(1981)006%3C0021:A0BIUF%3E2.0.CO;2.

LECUN, Y., BOTTOU, L., BENGIO, Y. et HAFFNER, P., 1998. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*. novembre 1998. Vol. 86, n° 11, pp. 2278-2324. DOI 10.1109/5.726791.

LIN, T.-Y., DOLLÁR, P., GIRSHICK, R., HE, K., HARIHARAN, B. et BELONGIE, S., 2017. *Feature pyramid networks for object detection* [en ligne]. 19 avril 2017. arXiv. arXiv:1612.03144. [Consulté le 5 août 2025]. Disponible à l'adresse : <http://arxiv.org/abs/1612.03144>

LIN, T.-Y., MAIRE, M., BELONGIE, S., BOURDEV, L., GIRSHICK, R., HAYS, J., PERONA, P., RAMANAN, D., ZITNICK, C. L. et DOLLÁR, P., 2015. *Microsoft COCO: Common Objects in Context* [en ligne]. 21 février 2015. arXiv. arXiv:1405.0312. [Consulté le 23 juin 2025]. Disponible à l'adresse : <http://arxiv.org/abs/1405.0312>

MERG, M.-L., DÉZERALD, O., KREUTZENBERGER, K., DEMSKI, S., REYJOL, Y., USSEGLIO-POLATERA, P. et BELLARD, J., 2020. Modeling diadromous fish loss from historical data: Identification of anthropogenic drivers and testing of mitigation scenarios. PINHEIRO, H. T. (éd.), *PLOS ONE*. 28 juillet 2020. Vol. 15, n° 7, pp. e0236575. DOI 10.1371/journal.pone.0236575.

NEYSHABUR, B., BHOJANAPALLI, S., MCALLESTER, D. et SREBRO, N., 2017. *Exploring generalization in deep learning* [en ligne]. 6 juillet 2017. arXiv. arXiv:1706.08947. [Consulté le 20 juin 2025]. Disponible à l'adresse : <http://arxiv.org/abs/1706.08947>

NYLUND, A., BJØRKNES, B. et WALLACE, C., 1991. Lepeophtheirus salmonis - A possible vector in the spread of diseases on salmonids. *Bulletin of the European Association of Fish Pathologists*. 1991. Vol. 11, n° 6, pp. 213-216.

RÄDSCH, T., REINKE, A., WERU, V., TIZABI, M. D., SCHRECK, N., KAVUR, A. E., PEKDEMIR, B., ROSS, T., KOPP-SCHNEIDER, A. et MAIER-HEIN, L., 2023. Labelling instructions matter in biomedical image analysis. *Nature Machine Intelligence*. 2 mars 2023. Vol. 5, n° 3, pp. 273-283. DOI 10.1038/s42256-023-00625-5.

REDMON, J., DIVVALA, S., GIRSHICK, R. et FARHADI, A., 2016. You Only Look Once: Unified, real-time object detection. In : *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* [en ligne]. Las Vegas, NV, USA : IEEE. juin 2016. pp. 779-788. [Consulté le 19 juin 2025]. ISBN 978-1-4673-8851-1. Disponible à l'adresse : <http://ieeexplore.ieee.org/document/7780460/>

REY, A. L., ASÍN, J., RUIZ ZARZUELA, I., LUJÁN, L., IREGUI, C. A. et DE BLAS, I., 2020. A proposal of standardization for histopathological lesions to characterize fish diseases. *Reviews in Aquaculture*. novembre 2020. Vol. 12, n° 4, pp. 2304-2315. DOI 10.1111/raq.12435.

SHALLU et MEHRA, R., 2018. Breast cancer histology images classification: Training from scratch or transfer learning? *ICT Express*. décembre 2018. Vol. 4, n° 4, pp. 247-254. DOI 10.1016/j.icte.2018.10.007.

TAYLOR, G. A., VOSS, S. D., MELVIN, P. R. et GRAHAM, D. A., 2011. Diagnostic errors in pediatric radiology. *Pediatric Radiology*. mars 2011. Vol. 41, n° 3, pp. 327-334. DOI 10.1007/s00247-010-1812-6.

TSENG, C.-H., HSIEH, C.-L. et KUO, Y.-F., 2020. Automatic measurement of the body length of harvested fish using convolutional neural networks. *Biosystems Engineering*. janvier 2020. Vol. 189, pp. 36-47. DOI 10.1016/j.biosystemseng.2019.11.002.

VON NEUMANN, John et MORGENTHAU, Oskar, 1944. *Theory of games and economic behavior*. Princeton University Press. United States. ISBN 978-0-691-13061-3.

W&B, 2025. Weights & Biases. *Weights & Biases* [en ligne]. 2025. [Consulté le 8 juillet 2025]. Disponible à l'adresse : <https://wandb.ai/site/>

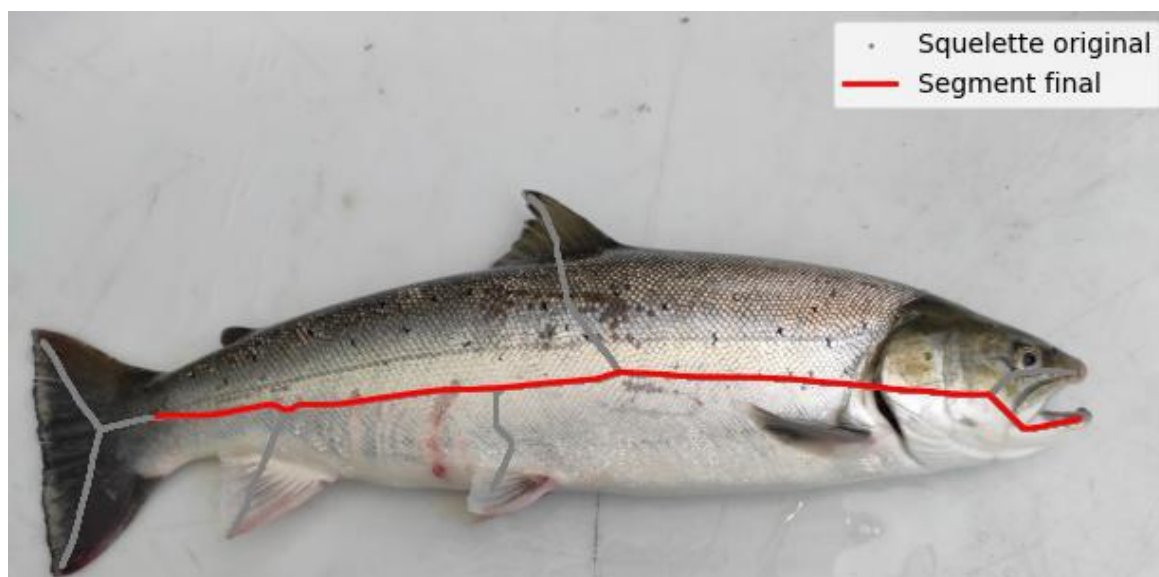
WEGENER, D. T. et PETTY, R. E., 1995. Flexible correction processes in social judgment: The role of naive theories in corrections for perceived bias. *Journal of Personality and Social Psychology*. 1995. Vol. 68, n° 1, pp. 36-51. DOI 10.1037/0022-3514.68.1.36.

YOUNGSON, A. F., BUCK, R. J. G., SIMPSON, T. H. et HAY, D. W., 1983. The autumn and spring emigrations of juvenile Atlantic salmon, *Salmo salar* L., from the Girnock Burn, Aberdeenshire, Scotland: environmental release of migration. *Journal of Fish Biology*. 1983. Vol. 23, n° 6, pp. 625-639. DOI 10.1111/j.1095-8649.1983.tb02942.x.

ZHANG, C., BRACKE, M., DA SILVA TORRES, R. et GANSEL, L. C.ian, 2024. Rapid detection of salmon louse larvae in seawater based on machine learning. *Aquaculture*. novembre 2024. Vol. 592, pp. 741252. DOI 10.1016/j.aquaculture.2024.741252.

Annexes

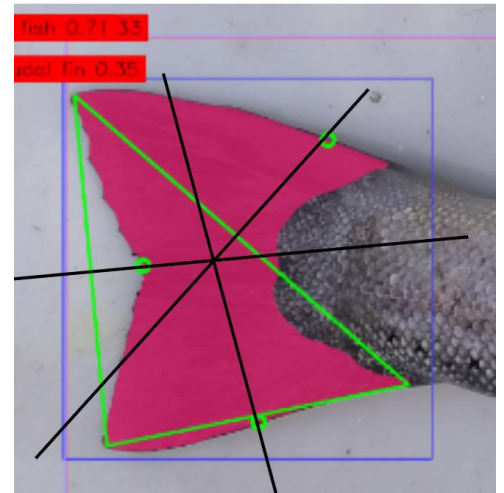
Annexe I : Squelette du masque du poisson



Annexe II : Démarche d'obtention du point fourche

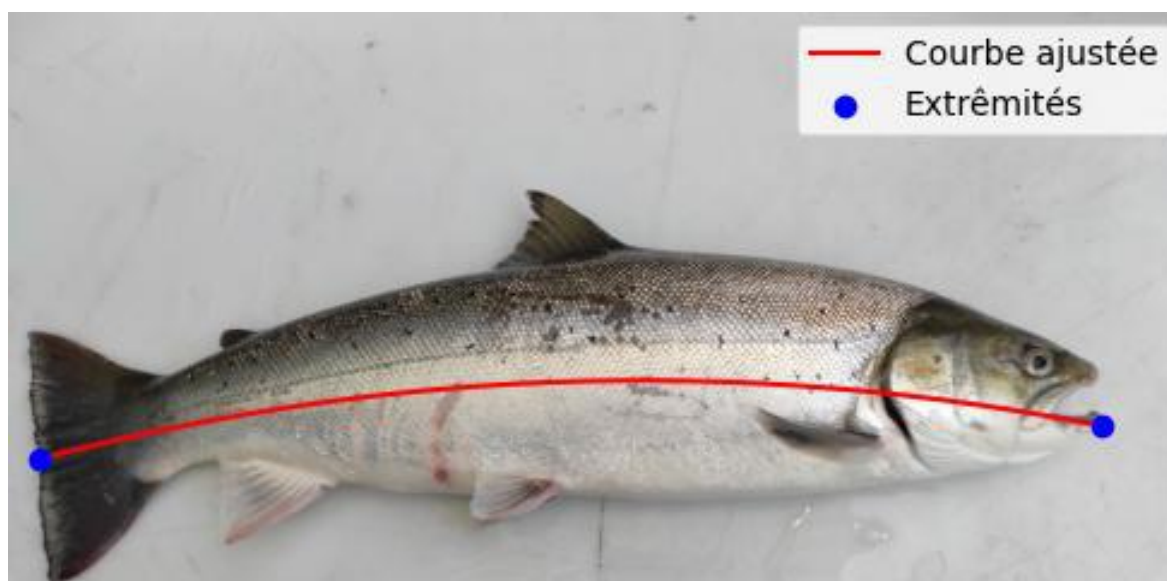
La procédure fut de :

- Extraire le contour du masque de la nageoire caudale
- Définir la plus grande longueur au sein de ce masque
- Définir un troisième point maximisant l'aire de masque dans le triangle formé par le point et le segment précédent comme base. Cette démarche permet d'extraire à coup sûr le segment entre les deux extrémités de la nageoire caudale (segment qui n'est pas systématiquement le plus long du masque, parfois inférieur à l'insertion de la nageoire jusqu'à une extrémité de celle-ci).
- Parcourir les médiatrices et conserver les points contenus dans le contour du masque, les plus proches du milieu de la base associée.

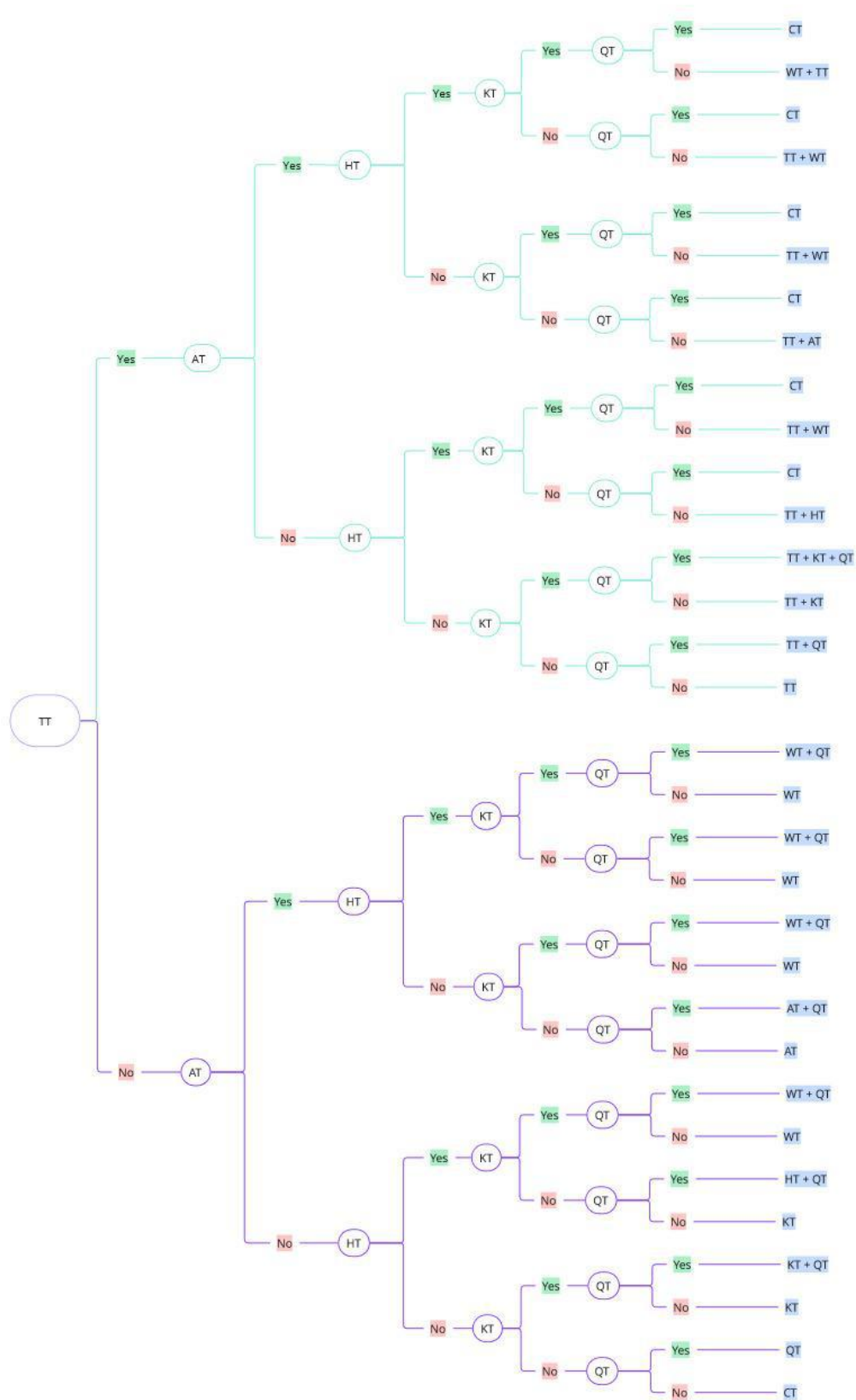


Cela nous donne trois points candidats à être le point fourche. Cette démarche géométrique est possible, car une fois le segment entre les extrémités de la nageoire caudale extrait, le point fourche se situe à la moitié de ce segment. Ceci explique le parcours des médiatrices pour trouver le point cible. Enfin, on sélectionne le point maximisant la distance avec l'extrémité de la bouche. Ces deux points nous permettent d'effectuer une régression de la courbure du squelette par un polynôme du second degré, dont le tracé est forcé de passer par les points définis (cf. Annexe III).

Annexe III: Courbe de régression du tracé du squelette

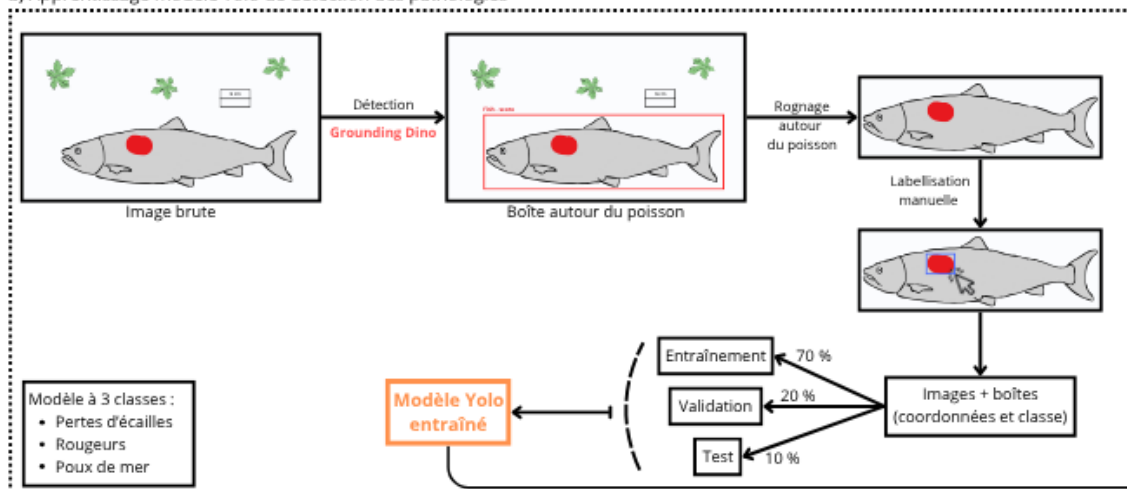


Annexe IV : Arbre de décision des codes de localisation

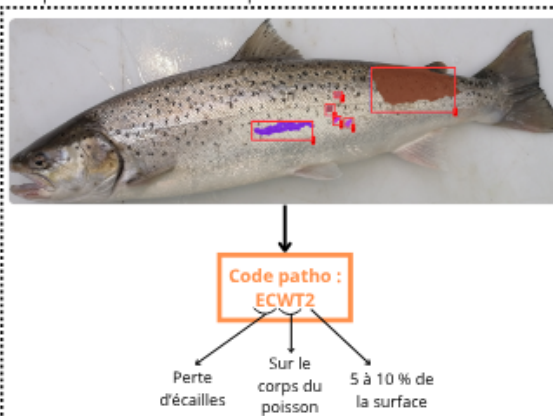


Annexe V : Schéma global de la démarche du stage : entraînement du modèle de détection des pathologies, pipeline et évaluation de la méthode

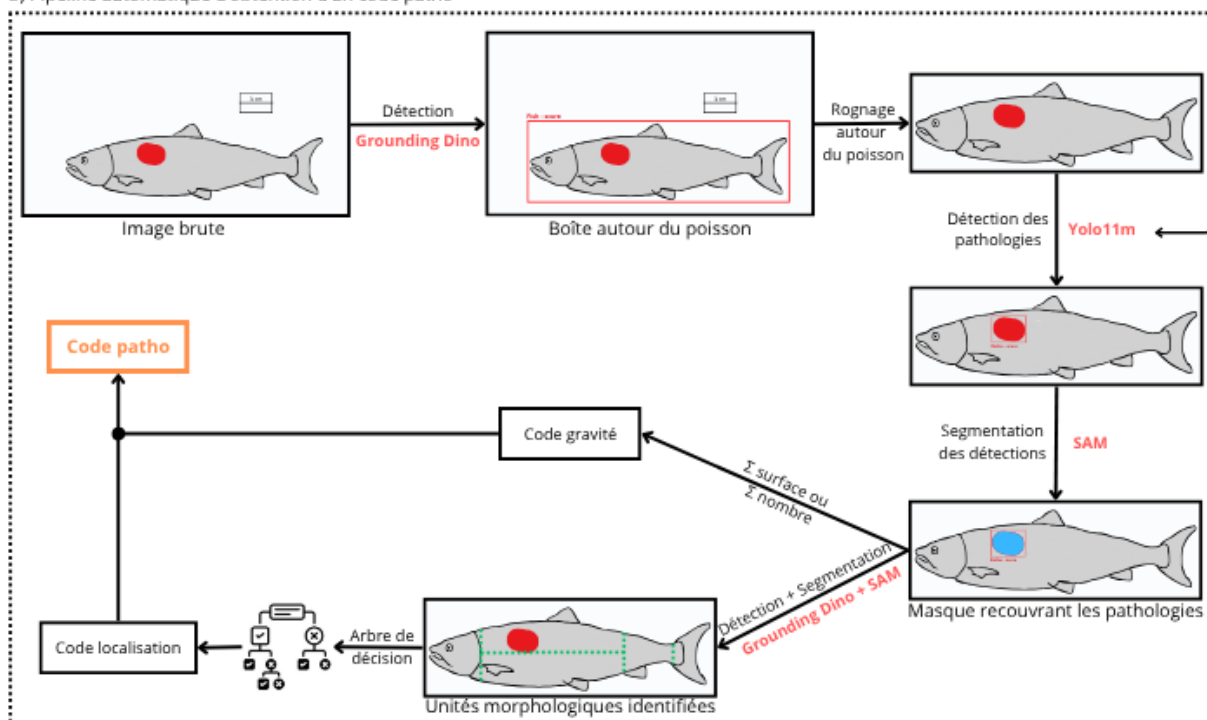
a) Apprentissage modèle Yolo de détection des pathologies



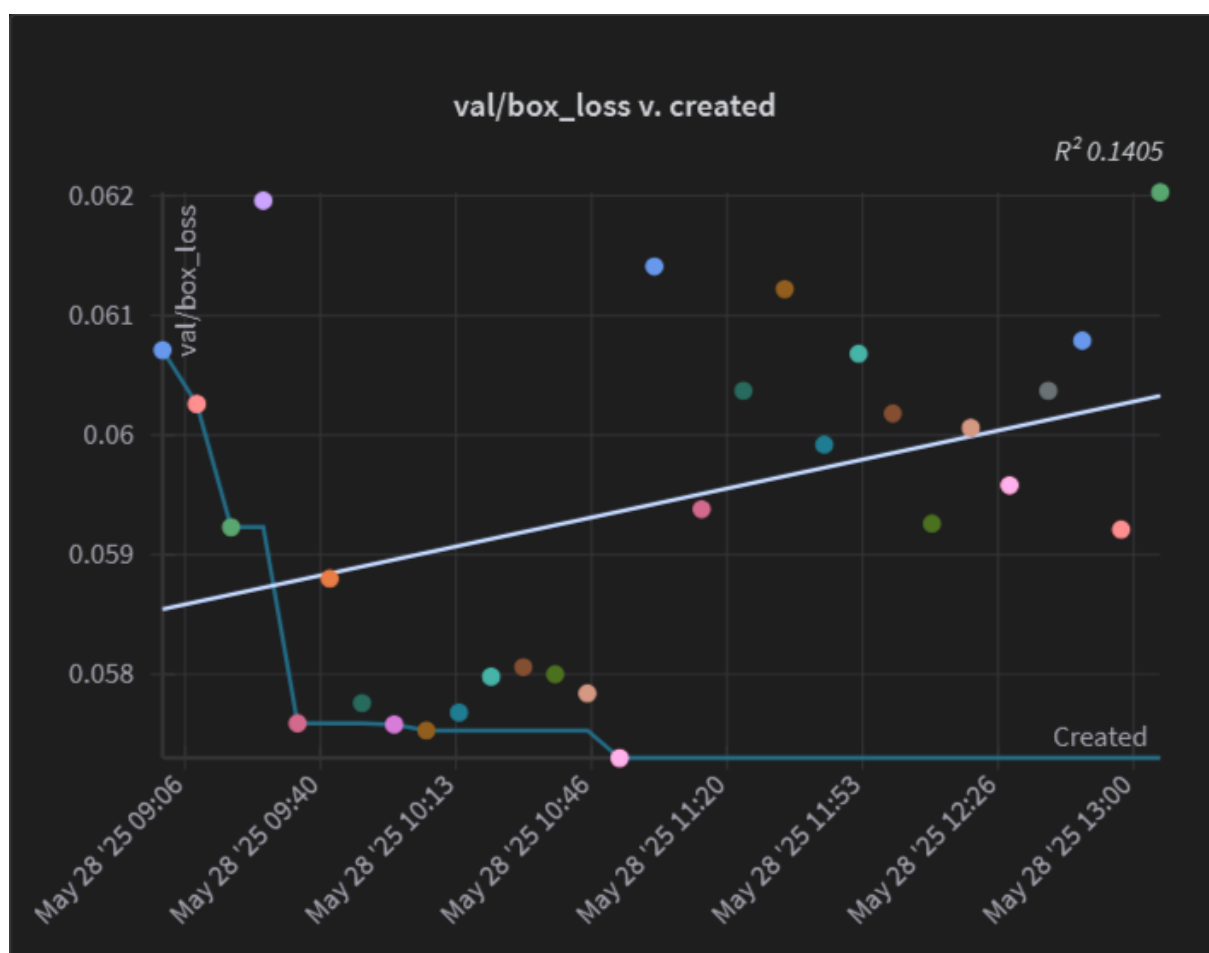
Exemple d'obtention du code patho



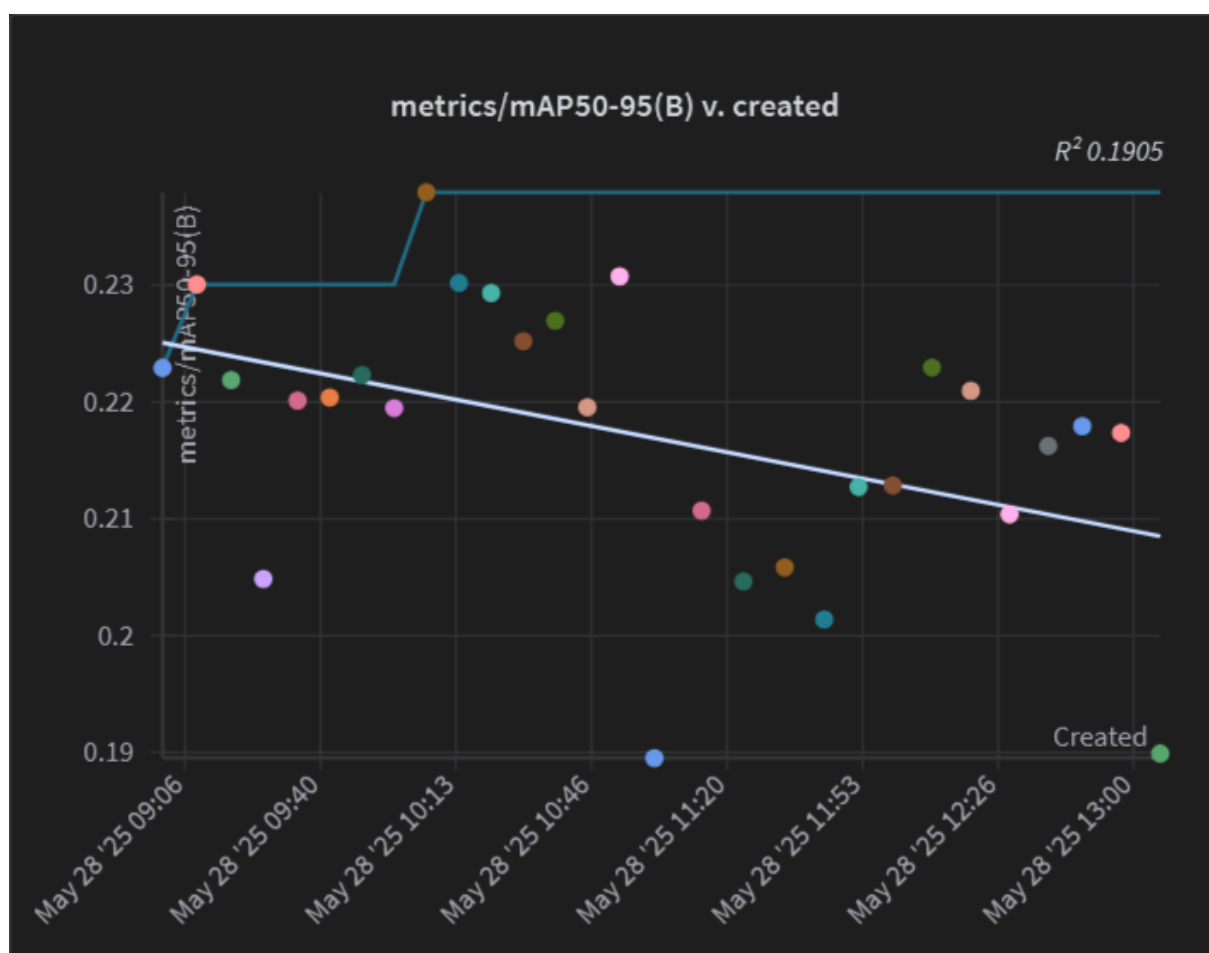
b) Pipeline automatique d'obtention d'un code patho



Annexe VI : Valeur de perte des prédictions des boîtes au cours des apprentissages dans le sweep



Annexe VII : Valeur de mAP50-95 au cours des apprentissages dans le sweep



	Diplôme : Ingénieur Spécialité : Ingénieur agronome Spécialisation / option : Sciences halieutiques et aquacoles Enseignant référent : Olivier Le Pape
Auteur(s) : Titouan Lecomte Date de naissance : 06/06/2001	Organisme d'accueil : UMR DECOD, INRAE & Institut Agro Rennes-Angers
Nb pages : 35 Annexe(s) : 7	Adresse : 64 Rue de Saint-Brieuc, 35042 Rennes
Année de soutenance : 2025	Maître de stage : François Martignac, Jérôme Guilton, Marie-Pierre Etienne
Titre français : Identification des pathologies sur photos de migrateurs amphihalins par analyse d'images Titre anglais : Identification of pathologies in pictures of anadromous fish through image analysis	
Résumé : Les poissons migrateurs amphihalins sont étudiés lors de leur passage au niveau des stations de contrôle des migrateurs, sur certains cours d'eau. Ce piégeage permet de recueillir des informations biométriques et de décrire les pathologies externes des poissons afin d'étudier ces espèces vulnérables. Ces dommages et parasites y sont décrits par des codes pathologie, compilant la nature, la localisation sur le poisson et la gravité de chaque pathologie. C'est pour répondre à une demande d'automatisation de la détection des dommages corporels que le projet PathologIA a été proposé. Il vise à appliquer des outils d'analyse d'images par apprentissage profond. L'utilisation de modèles pré-entraînés et le développement de modèles à apprentissage supervisé ont permis d'extraire et traiter les informations contenues sur les photographies. Le pipeline mis en place est conçu pour détecter les pathologies, séparer le poisson en différentes unités morphologiques et ainsi caractériser la localisation et la gravité de chaque pathologie identifiée. Selon que l'on cherche à maximiser l'exhaustivité ou la précaution du modèle à prédire une absence ou une présence de chaque pathologie, nous avons sélectionnés des seuils de confiance propres à chaque pathologie. Cette méthode a permis l'estimation de codes patho comparables à ceux décrits par les opérateurs sur le terrain. Les limites des modèles pourront être compensées par un nombre d'images plus important dans la base de données utilisée, notamment pour l'entraînement du modèle de détection.	
Abstract: Anadromous migratory fish species are monitored during their passage through fish migration control stations, on rivers. This trapping process allows the collection of biometric data and the description of external pathologies observed on the fish's body, in order to study population's health trends of these vulnerable species. These diseases and parasites are recorded using pathology codes, which compile the nature, location on the body, and severity of each damage. To meet the need for automating the detection of physical damage, the PathologIA project aims to apply deep learning-based image analysis tools. The use of pre-trained models, along with supervised training models, has enabled the extraction and processing of information from photographs. The established pipeline is designed to detect pathologies, to segment the fish's body into different morphological units, and characterize the location and severity of each identified pathology. Depending on whether the goal is to maximize comprehensiveness or caution in predicting the presence or absence of a pathology, a threshold must be selected for each pathology. This method enables the estimation of pathology codes with promising and useable quality for field operators. Any lack of precision could potentially be offset by supplementing the image database and labels used for the training process of the model.	
Mots-clés : Apprentissage profond, Analyse d'images, Amphihalin, Salmonidés, Pathologie Key Words: Deep learning, Image analysis, Anadromous, Salmonids, Pathology	