

Chapitre 4

État de l'art

Examinons maintenant les diverses techniques de tatouage des objets 3D. Elles empruntent des éléments déjà validés en 2D. Par exemple, on cherchera à répéter la marque et à enfouir un même bit de nombreuses fois de façon à compenser la fragilité du tatouage par une forte redondance. On utilisera également souvent des clefs secrètes afin de brouiller soit l'ordre de parcours des diverses facettes, soit l'ordre d'apparition des bits de la marque. Néanmoins, comme nous allons le voir, la géométrie des objets maillés offre de nombreuses solutions originales aux tatoueurs.

Nous classons les différents schémas suivant qu'ils font ou non appel à une configuration géométrique pour insérer la marque. En effet, les techniques fondées sur une configuration géométrique reposent toutes sur le même paradigme d'*arrangement* de la marque. A l'inverse, les autres méthodes ont pour point commun de disperser l'information utile de manière plus fine sur le maillage.

Avant d'aller plus loin, il convient d'avertir le lecteur que les références en tatouage 3D ne détaillent que trop rarement de façon précise la manière dont l'information est exactement encodée. Cet état de l'art est donc aussi précis que les références nous l'ont permis. En particulier, la description des schémas par invariants géométriques reste particulièrement lapidaire sur le sujet (quand on ne se contente pas de mentionner uniquement la nature de l'invariant que l'on modifie). Pour nous, nous tenterons au contraire au chapitre 5 de détailler le mieux possible notre méthode d'insertion géométrique.

4.1 Schémas avec configuration géométrique locale

L'information de connexité peut être mise à profit dans le cadre du tatouage. Les schémas utilisant cette information font appel à une configuration géométrique : un ou plusieurs triangles voisins. L'information cachée est alors codée en modifiant des invariants géométriques. Du type d'invariant retenu dépend la robustesse *a priori* du schéma, on choisit donc les invariants en fonction des transformations que subiront les objets [27]. Nous en donnons une liste non exhaustive :

- Sont invariants par rotation et translation : la longueur d'un segment, l'aire d'un polygone et le volume d'un polyèdre.
- Sont invariants par transformée affine : le rapport des longueurs de deux segments

colinéaires, le rapport de l'aire de deux polygones coplanaires et le rapport des volumes de deux polyèdres.

- Seul le birapport de quatre points alignés est invariant par transformation projective [25].

Les méthodes retenues vont alors introduire de légères modifications dans les invariants choisis. Comme il n'y a pas d'invariant, à part la valence, aux perturbations des positions des sommets, elles sont sensibles à l'ajout de bruit. Elles pourront résister à des modifications locales de la connexité pourvu que la localisation soit bonne. On distingue trois manières de localiser l'information utile, appelées arrangements.

- Soit on utilise un parcours canonique du graphe de connexité (comme dans les méthodes de compression), l'arrangement est dit global,
- Soit on utilise un parcours canonique de chacune des sous-parties du maillage, et l'arrangement est dit local,
- Soit enfin on marque explicitement la *localisation* de l'information. On est alors confronté à deux exigences : marquer une configuration géométrique pour éviter de l'utiliser une seconde fois, et aussi inscrire un index permettant de replacer l'information cachée au sein de la marque. Il faut, lors de la détection, savoir distinguer les configurations géométriques codant l'information utile des autres, et il faut savoir quel bit du message telle configuration code effectivement. L'arrangement est alors dit indexé.

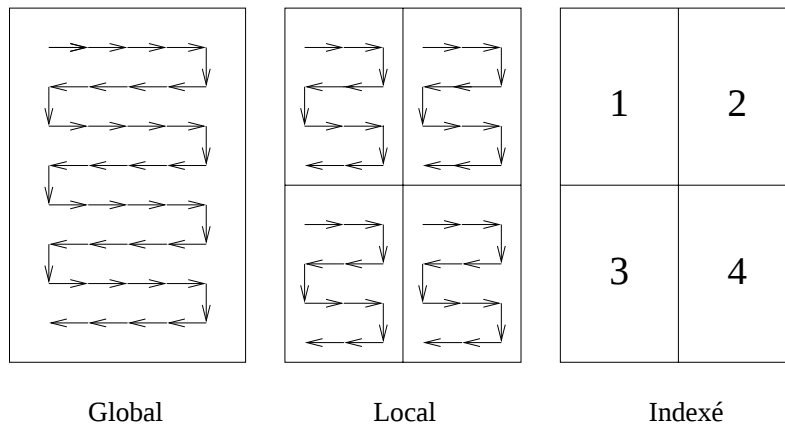


FIG. 4.1 – Les trois types d'arrangements disponibles. Les arrangements globaux et locaux nécessitent un parcours du maillage : il faut alors choisir un germe relativement stable. Seul l'arrangement indexé permet une relecture en aveugle sans nécessiter de parcours particulier du graphe de connexité.

Les deux premiers types d'arrangements doivent trouver un germe⁷⁶ pour initialiser le parcours canonique du graphe de connexité. On peut utiliser deux techniques pour cela :

- des séquences de synchronisation, qui sont des marques connues qui débutent le message utile et que l'on retrouve par des séquences d'essai-erreur,

⁷⁶Le germe d'un parcours est soit un triangle, une arête ou un sommet. Le type du germe dépend des spécificités du parcours. Le choix du germe détermine le parcours, un peu comme la graine d'un générateur pseudo-aléatoire détermine la suite de nombre fournis par le générateur.

- des initialisations sur des configurations remarquables du volume ou de la partition.

On peut par exemple utiliser comme germe, s'il s'agit d'un côté, celui dont le volume du tétraèdre sous-tendu par ses deux triangles incidents est maximal. Comme les rapports des volumes de polyèdres restent inchangés dans le cas de transformations affines, cette initialisation sera robuste à ces transformations. De même, s'il fallait un germe qui soit un triangle, le triangle d'aire la plus grande ferait l'affaire dans le cas de transformations rigides (rotation, translation, etc.) mais il n'est pas robuste à une transformation affine.

4.1.1 Arrangement global

Dans [8], une méthode de tatouage (appelée *triangle flood*) fondée sur un parcours global du graphe de connexité est présentée. Ce parcours peut être vu sous forme d'arbre. Partant d'un triangle, on vient collecter les sommets qui forment un triangle avec un côté adjacent du triangle initial. On s'intéresse aux hauteurs issues de ces sommets dans les triangles adjacents. On ordonne ensuite ces hauteurs, à la fois pour fixer l'ordre de parcours des triangles, et aussi pour insérer l'information à cacher. On itère ensuite sur les triangles adjacents, en parcourant l'arbre *en largeur d'abord*.

C'est l'ordre des hauteurs qui fixe l'ordre des bits de la marque. On impose que les hauteurs, ordonnées, aient une distance d_{min} entre chacune d'elles successivement. La première modification des hauteurs a pour but d'imposer la contrainte sur d_{min} . Le décodeur peut alors reconstituer le parcours en réordonnant les hauteurs, et en ne s'occupant que de celles dont les écarts successifs sont supérieurs à d_{min} . La deuxième opération vise à inscrire les valeurs de bits de la marque, en modifiant encore les hauteurs. A l'inscription de la marque, on respecte la contrainte sur d_{min} afin de ne pas effacer le parcours que l'on vient de créer (nous retrouverons ce problème de causalité au chapitre 5).

L'intervalle admissible de modification des hauteurs est calculé de manière à :

- Conserver le parcours (contrainte à respecter sur d_{min})
- Insérer m bits par modification de hauteur (l'auteur prend $m = 2$)

On a représenté un maillage et l'arbre correspondant sur la Fig. 4.2. On notera que l'on peut gérer des maillages qui ne sont pas nécessairement des variétés (voir les sommets 4 et 5), et que le cas spécial où un même sommet appartient à deux triangles adjacents au triangle père dans l'arbre génère une exception (considérer le sommet 2).

Un tel parcours n'admet aucune modification de la connexité et instaure un ordre global d'arrangement des bits de la marque sur le maillage. Le décodage prend fin lorsque tous les bits sont retrouvés (la longueur du message caché est supposée connue au détecteur), ou que l'on détecte une séquence spéciale de bits indiquant la fin du message (solution non retenue par l'auteur).

4.1.2 Arrangement local

Nous donnons ici deux exemples d'arrangement local de l'information cachée. Le premier se sert de sites admissibles proches pour déterminer le numéro du bit que l'on va cacher. Le second découle quant à lui d'un arrangement global : là encore la marque est enfouie selon un arbre de recouvrement des sommets, mais les auteurs partitionnent le maillage au

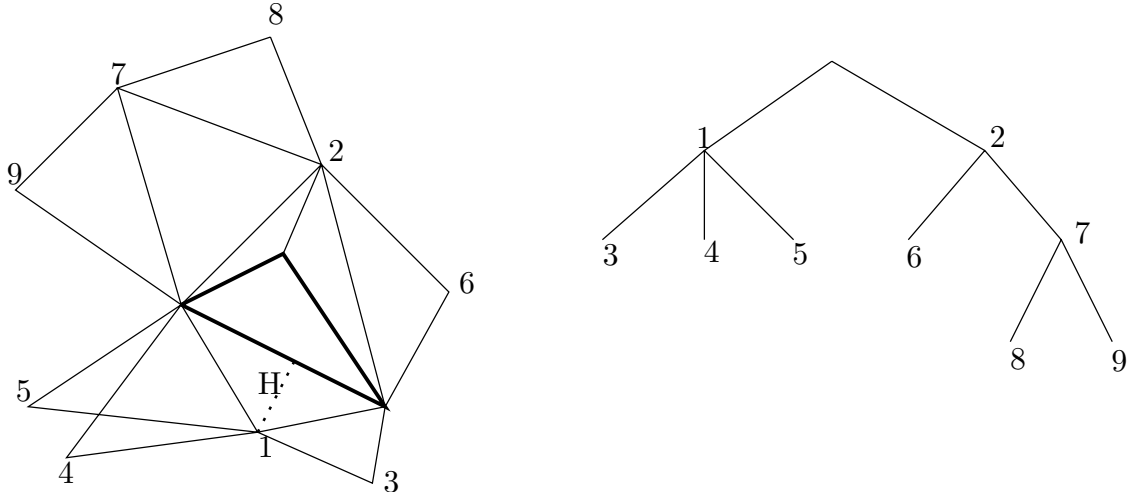


FIG. 4.2 – Parcours global du maillage dans l'algorithme *triangle flood* [8] en largeur d'abord. Le triangle initial est représenté en gras à droite, l'arbre obtenu est à droite. Le sommet 2 génère une exception (il est commun à deux triangles adjacents au triangle en gras) de manière à ce que l'arbre ne le contienne qu'une seule fois. Ce parcours permet de gérer des maillages qui ne sont pas nécessairement des variétés (les sommets 4 et 5 partagent une arête commune à trois triangles).

préalable et appliquent un arrangement global sur chacune des partitions. En partitionnant le maillage, on peut rendre local un arrangement global.

Schéma de Harte

Dans [35], la marque est répétée sur des parties du maillage qui ne se recouvrent pas. Pour cela, une sélection des sommets pouvant être modifiés est dressée. On considère qu'un sommet peut être marqué s'il n'appartient pas à un segment trop long. On calcule alors la somme D_i des longueurs de chaque sommet p_i à ses voisins :

$$D_i = \sum_{p_j \in p_i^*} \|p_i p_j\|. \quad (4.1)$$

Un sommet p_i est déclaré valide si pour un certain seuil T :

$$D_i < T. \quad (4.2)$$

Tous les sommets voisins sont déclarés invalides. L'idée de [35] consiste à se servir du voisinage d'un sommet valide p_i pour établir un espace de tatouage dans lequel on modifie p_i . On remarque que les voisins de p_i valides peuvent définir un ellipsoïde (voir Fig. 4.3) centré autour de μ_i :

$$\mu_i = \frac{1}{d_{p_i}} \sum_{p_j \in p_i^*} p_j. \quad (4.3)$$

La forme de l'ellipsoïde est donnée par la matrice de covariance S_i :

$$S_i = \frac{1}{d_{p_i}} \sum_{p_j \in p_i^*} (p_j - \mu_i)(p_j - \mu_i)^T. \quad (4.4)$$

La surface de l'ellipsoïde est constituée des points p tels que :

$$(p - \mu_i)^T S_i^{-1} (p - \mu_i) = K, \quad (4.5)$$

où K est un facteur de normalisation. Soit M_t le nombre de sommets valides pour l'insertion de B bits. Les M_t sommets sélectionnés sont partagés aléatoirement en plusieurs ensembles de B sommets. A l'intérieur de chacun de ces ensembles, on vient ordonner les B sommets qui le composent en se fondant sur les distances entre les centres de leur ellipsoïde. Associés aux numéros de 1 à B , les sommets sont donnés par :

$$i = \arg \min_j \sum_{k=1, k \neq j}^{B-1} \|\mu_k - \mu_j\|^2, \quad (4.6)$$

où l'on ne considère que les $B - 1$ plus proches ellipsoïdes pour le calcul. Une fois qu'un sommet a obtenu son numéro, son ellipsoïde ne rentre plus en ligne de compte pour l'attribution des autres numéros.

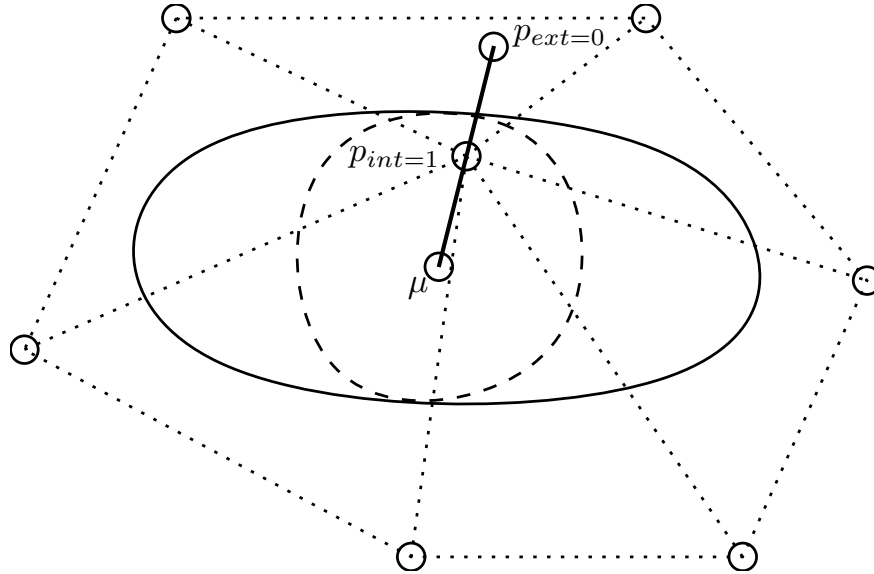


FIG. 4.3 – L'espace de tatouage associé à l'ellipsoïde dans [35]. Ici p_i est situé à l'intérieur de l'ellipsoïde et code un 1 par convention. Si l'on voulait cacher un 0, il faut amener p_i en p_{ext}

Une fois les numéros attribués aux sommets, il ne reste plus qu'à insérer l'information. Ici, l'ellipsoïde matérialise un intérieur et un extérieur. L'intérieur sert à coder les 1, l'extérieur les 0. Si le sommet code déjà dans la bonne valeur de bit, il n'est pas modifié, sinon il est poussé de l'intérieur vers l'extérieur de l'ellipsoïde (pour passer de 1 à 0), ou l'inverse (pour passer de 0 à 1). Le déplacement se fait selon la droite $\mu_i p_i$.

Lorsque les B premiers bits sont insérés, on renouvelle le même algorithme pour les B sommets de l'ensemble suivant de sommets valides. Chaque bit est donc répété $\lfloor M_t/B \rfloor$ fois. L'ordre des bits est uniquement fonction des $B - 1$ plus proches ellipsoïdes, l'arrangement est donc local.

Schéma TVR

L'invariant utilisé dans le schéma TVR [56] est le rapport des volumes de deux tétraèdres. Pour cette raison, ce schéma nécessite un arrangement global ou local (version TVR Cluster). Il faut en effet un tétraèdre de référence auquel rapporter le volume des autres. Les tétraèdres sont formés par une arête (p_3p_4 sur la figure 4.4) et les deux triangles incidents ($p_1p_3p_4$ et $p_2p_3p_4$). Le parcours du graphe de connexité repose sur un arbre de recouvrement des points.

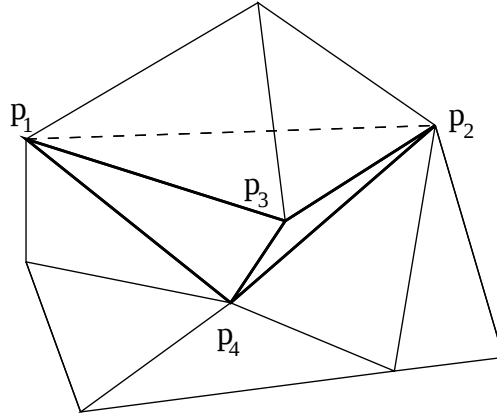


FIG. 4.4 – Schéma *Tetrahedral Volume Ratio* : encodage du parcours unique du maillage. Tous les volumes des tétraèdres sont rapportés à celui du tétraèdre initial (le premier dans le parcours du graphe). On insère l'information en modifiant les coordonnées des points des tétraèdres au numérateur (dans les rapports de volumes).

Une fois cet arbre construit, on se sert des arêtes qu'il définit pour référencer les tétraèdres et leur ordre de parcours. Pour retrouver le début du parcours canonique (global ou local), on procède par essais successifs jusqu'à trouver une séquence de synchronisation réputée connue. Dans le cas d'un arrangement local, les sous-parties du maillages sont définies par les branches de l'arbre de recouvrement des points que l'on utilise.

4.1.3 Arrangement indexé

Le schéma TSQ [56] est un des deux seuls exemples connus d'implantation d'un arrangement indexé avec [82] (nous réaliserons notre propre arrangement indexé au chapitre 5). Ce schéma s'appuie sur une configuration géométrique composée de quatre triangles, chaque configuration contenant deux bits d'information cachée. La capacité d'insertion est donc la moitié de celle du schéma précédent. L'insertion repose sur la modification de paires de rapports entre les bases des triangles et les hauteurs associées. La configuration est constituée d'un triangle, central, et de ses trois triangles adjacents. Le triangle central valide la présence

d'information cachée dans la configuration, un des trois triangles adjacents sert à coder le numéro du bit caché, et les deux derniers triangles codent l'information utile. Les auteurs ne précisent pas davantage comment ils codent géométriquement l'information, même si l'on peut supposer qu'ils cherchent à quantifier des rapports de grandeurs géométriques.

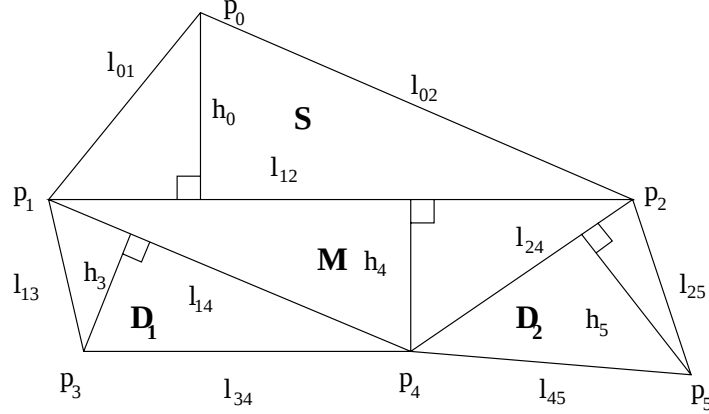


FIG. 4.5 – Schéma *Triangle Similarity Quadruple* : la configuration géométrique est composée de quatre triangles adjacents. Le triangle central valide la configuration comme porteuse d'information cachée. Le numéro du bit caché et deux valeurs indexées sur le numéro sont cachées dans les trois triangles restants.

4.2 Schémas sans configuration géométrique locale

Afin de s'affranchir de la contrainte de connexité à base de triangles, certains schémas de tatouage se basent sur des points uniquement. Les calculs de tatouage peuvent faire intervenir le voisinage local du sommet, mais la relecture ne nécessitera pas de retrouver des configurations géométriques locales dans lesquelles calculer des invariants. Ces schémas utilisent en général une référence à l'objet entier et se révèlent alors sensibles à l'attaque de coupe, bien que leur robustesse face à d'autres attaques plus complexes comme le remaillage ou la simplification soit élevée. Les deux premières approches présentées ici cherchent soit à modifier la courbure locale de l'objet, soit à modifier la distance des points au centre de masse de l'objet.

La troisième approche vise au contraire à renseigner l'utilisateur sur l'intégrité du maillage. Le marquage devra permettre de vérifier un certain calcul pour tous les sommets. On utilisera alors un autre espace non pour davantage de robustesse mais pour davantage de sécurité.

4.2.1 Schéma de Wagner

Ce schéma [79] insère l'information dans les valeurs relatives des normales locales en chaque sommet : il modifie donc la courbure locale de l'objet. Les normales locales sont

calculées ainsi :

$$n_i = \frac{1}{d_{p_i}} \sum_{v \in p_i^*} (v - p_i). \quad (4.7)$$

On calcule la longueur moyenne de ces normales n_i sur l'ensemble du maillage :

$$d = \frac{1}{L} \sum_{i=0}^L \|n_i\|. \quad (4.8)$$

Et l'on quantifie les normales sur c/d niveaux, c étant une clef secrète. On définit alors le laplacien quantifié k_i par :

$$k_i = \left\lfloor \frac{c}{d} \|n_i\| \right\rfloor. \quad (4.9)$$

La marque est calculée à partir d'une fonction $f(p)$ continue définie sur la sphère unité. La marque f peut alors être un logo projeté sur la sphère unité, ou une séquence pseudo-aléatoire. De la même manière que ci-dessus, on quantifie la valeur que prend f dans chaque direction n_i en valeurs w_i :

$$w_i = \left\lfloor 2^b f\left(\frac{n_i}{\|n_i\|}\right) \right\rfloor, \quad (4.10)$$

où b est le nombre de bits servant à coder chaque w_i . L'insertion remplace b bits de k_i par ceux de w_i , on obtient alors les entiers k'_i . On calcule les vecteurs laplaciens modifiés :

$$n'_i = \frac{k'_i d}{c} \frac{n_i}{\|n_i\|}, \quad (4.11)$$

et on détermine les nouvelles coordonnées p'_i et v' des points qui satisfont les équations suivantes :

$$n'_i = \frac{1}{d_{p'_i}} \sum_{v' \in p'_i} (v' - p'_i). \quad (4.12)$$

Ce système de $L + 1$ équations est singulier : on ne peut reconstruire un objet uniquement à partir de ses normales locales. On laisse alors environ 20% des points inchangés. Puisque la majorité des points ont été modifiés, on doit calculer le nouveau paramètre global c' pour la relecture en fonction de la nouvelle moyenne d_n des normales locales modifiées :

$$c' = c \frac{d}{d_n} \quad (4.13)$$

Au lieu de la norme euclidienne classique, on peut choisir d'utiliser une autre norme, invariante par transformation affine, telle que décrite dans [55].

Pour la relecture, on cherche à retrouver la fonction f qui fait office de marque. Toutefois, l'auteur n'explique pas comment, même s'il l'affirme, il peut rendre son tatouage indécélable, ni avec quelle sécurité.

4.2.2 Schéma *Vertex Flood*

L'algorithme *Vertex Flood* [7, 8] repose sur une modification de la distance des points au centre de gravité de l'objet. On choisit une distance maximale D_{max} au centre de gravité de l'objet. L'intervalle $[0; D_{max}]$ est alors subdivisé en $m = 2^N$ intervalles : c'est l'espace de tatouage. Les sommets sont déplacés sur la droite passant par le centre de gravité et le sommet à tatouer, pour amener le sommet à la distance voulue du centre de gravité. On peut cacher plusieurs marques en sélectionnant différents sous-ensembles de points pour chacune d'entre elles. Toutefois, une coupe modifiant l'ensemble des points sur lesquels on calcule le centre de gravité, la marque est perdue par une coupe. En revanche, ce schéma s'adapte naturellement aux maillages quelconques (soupe de polygones, maillages non nécessairement conformes, dégénérescences, entre autres).

4.2.3 Schéma pour l'intégrité

Dans [82], Yeo et Yeung proposent une méthode permettant d'attester de l'intégrité de la position de chaque sommet. Comme souvent avec les schémas pour l'intégrité, il vaut mieux expliquer d'abord comment fonctionne le décodeur. Soit p_i un sommet dont on souhaite tester l'intégrité. On va chercher à générer deux valeurs à partir de p_i : l'index de *localisation* $L(p_i)$ et l'index de *valeur* $v(p_i)$. On utilise $L(p_i)$ pour fixer l'index du bit de la marque W que l'on est en train de relire. Un bit du tatouage est repéré par sa position dans un tableau grâce à l'index de localisation, et l'index de valeur code effectivement sa valeur. Un sommet du maillage va coder un bit du tatouage. Parallèlement, une clef de vérification K , indexée par $v(p_i)$, permet d'attester que le sommet est intègre si :

$$K(v(p_i)) = W(L(p_i)) \quad (4.14)$$

Si les valeurs diffèrent, le sommet est déclaré non intègre. Le but du tatouage est ici de faire en sorte que tous les sommets soient intègres au regard de l'Eq. 4.14. Une mauvaise clef K conduira à déclarer seulement un sommet sur deux intègre en moyenne.

Nous commençons par expliquer comment calculer l'index de localisation. La marque W est ici une image binaire de taille 64×64 . L'index de localisation $L(p_i)$ sera en réalité composé de deux index : $L_x(p_i)$ en abscisse et $L_y(p_i)$ en ordonnée. Pour les calculer, on commence par définir l'ensemble $\mathcal{N}(p_i)$ constitué de p_i et de ses voisins $p_j \in p_i^*$ tels que $j < i$. On peut ainsi modifier les sommets dans l'ordre dans lequel ils sont énumérés sans risquer d'effacer ce qu'on vient d'écrire (problème de causalité). Il faut ensuite s'intéresser au barycentre s de p et de ses voisins :

$$s = \frac{1}{1 + d_p} \sum_{p_i \in p \cup p^*} p_i. \quad (4.15)$$

Ensuite, muni d'une fonction q de normalisation et de quantification, on établit les index de localisation à l'aide de deux fonctions f_x et f_y :

$$L_x(p) = f_x(q(s_x), q(s_y), q(s_z)), \quad (4.16)$$

$$L_y(p) = f_y(q(s_x), q(s_y), q(s_z)), \quad (4.17)$$

où f_x et f_y servent à assigner un couple d'index à p (les auteurs semblent suggérer une combinaison des $q(s)$ modulo la taille de la marque). On effectue ainsi l'équivalent d'une paramétrisation du barycentre s .

Pour calculer l'index de valeur, on part de p et, conservant la fonction q , on a besoin de trois tables secrètes K_1 , K_2 et K_3 combinées entre elles par un ou exclusif :

$$K(v(p)) = K_1(q(p_x)) \oplus K_2(q(p_y)) \oplus K_3(q(p_z)). \quad (4.18)$$

Ces tables mesurent 256 bits chacune. Elles peuvent découler de séquences de bits issues d'un générateur pseudo-aléatoire pour plus de souplesse. Du point de vue de la sécurité, la taille de l'espace d'attaque exhaustive sera donc de $256 \times 256 \times 256 = 2^{24}$ clefs.

Enfin, le décodeur saura fournir la proportion c de sommets non entiers :

$$c = \frac{\#\{p | K(v(p)) \neq W(L(p))\}}{N}. \quad (4.19)$$

Si le maillage n'a subi aucune modification, on a naturellement $c = 1$.

Le but du tatouage est donc de perturber localement chaque sommet de manière à ce que l'Eq. 4.14 soit vérifiée en chacun d'eux. Pour insérer l'information, on choisit de perturber p en p' jusqu'à ce que $K(v(p'))$ prenne la valeur désirée. Mais il faut prêter spécialement attention au fait qu'en choisissant une telle stratégie on se doit de respecter un ordre sur le parcours des sommets qui soit le même au codeur et au décodeur. Dans [82], on utilise l'ordre de description des sommets dans le fichier représentant le maillage. C'est là le principal écueil de cette méthode, qui ne tolère donc aucune modification de la connexité, consisterait-elle seulement en une renumérotation des sommets.

Afin de perturber p en p' , on met en œuvre un algorithme itératif. On choisit des pas de perturbation Δ_x , Δ_y et Δ_z que l'on incrémente doucement d'une itération à l'autre. Le début de l'exécution de l'algorithme est donné par :

$$\begin{aligned} p' &\leftarrow (p_x + \Delta_x, p_y, p_z), \\ p' &\leftarrow (p_x - \Delta_x, p_y, p_z), \\ p' &\leftarrow (p_x, p_y + \Delta_y, p_z), \\ p' &\leftarrow (p_x, p_y - \Delta_y, p_z), \\ p' &\leftarrow (p_x, p_y, p_z + \Delta_z), \\ p' &\leftarrow (p_x, p_y, p_z - \Delta_z), \\ p' &\leftarrow (p_x + \Delta_x, p_y + \Delta_y, p_z), \\ &\text{etc...} \end{aligned}$$

Après relecture de la marque, l'utilisateur dispose de la valeur c pour estimer l'ampleur des dégradations subies par le maillage. Eventuellement, on pourra faire aussi apparaître dans un visualiseur les sommets non entiers. Enfin, une solution à mi-chemin consiste à prendre un logo bien défini pour marque W . Si le maillage n'a subi aucune modification, l'image extraite sera parfaitement nette. A l'inverse, plus les modifications auront porté sur un nombre important de sommets, plus l'image W apparaîtra bruitée. Cette solution peut être envisagée si l'inspection visuelle en détails des sommets non entiers est trop coûteuse.

4.3 Autres espaces de représentation

Nous présentons ici deux schémas de tatouage qui utilisent un espace différent de représentation des données maillées 3D. Le premier tire parti de la décomposition spectrale, qui sera abordé au chapitre 6. Le deuxième se place dans la représentation multirésolution *Progressive meshes* due à Hoppe [37]. Nous présenterons encore un troisième schéma, spécifiquement destiné à la CAO, qui permet de tatouer des NURBS.

4.3.1 Schéma dans l'espace spectral de la géométrie

Comme nous le présenterons en détails au chapitre 6, on peut décomposer la géométrie du maillage sur une base de fonctions propres. On effectue la projection des trois vecteurs de coordonnées X , Y et Z pour obtenir leurs transformés P , Q et R . L'espace de projection est dit spectral, et les transformés P , Q et R sont composés de coefficients spectraux. Cet espace a pu être utilisé pour le tatouage [58]. Ce schéma propose une méthode additive d'insertion de la marque dans le domaine spectral de décomposition de la géométrie. Il nécessite une resynchronisation géométrique sur le modèle original préalable à la relecture (imposé par les propriétés de non-invariance en rotation de la décomposition spectrale). Par répétition du message c fois, on obtient la séquence b :

$$b_i = m_j \quad , \quad jc \leq i < (j+1)c. \quad (4.20)$$

Ensuite on vient moduler classiquement la séquence des b_i :

$$b'_i = \begin{cases} -1 & \text{si } b_i = 0 \\ 1 & \text{sinon.} \end{cases} \quad (4.21)$$

L'insertion se fait alors additivement dans l'espace transformé de la géométrie à l'aide d'une séquence pseudo-aléatoire s binaire à valeur dans $-1, 1$:

$$\begin{cases} P'_i = P_i + \alpha \cdot b'_i \cdot s_i \\ Q'_i = Q_i + \alpha \cdot b'_i \cdot s_i \\ R'_i = R_i + \alpha \cdot b'_i \cdot s_i, \end{cases} \quad (4.22)$$

où α représente la force d'insertion. L'extraction de la marque repose sur un alignement préalable entre le maillage tatoué et l'original. Cet alignement opéré, la décomposition spectrale du maillage dont on souhaite extraire la marque donne les coefficients $\hat{P}$, $\hat{Q}$, $\hat{R}$, et celle du maillage original les coefficients P , Q , R . On calcule ensuite la corrélation globale q_j :

$$q_j = \frac{1}{3} \sum_{X \in \{P, Q, R\}} \sum_{i=j.c}^{(j+1).c-1} (\hat{X}_i - X_i) \cdot s_i. \quad (4.23)$$

Pour obtenir la séquence estimée $\hat{b}'$, il suffit de tester le signe de chaque élément q_i :

$$\hat{b}'_i = \text{signe}(q_i). \quad (4.24)$$

La démodulation pour obtenir la séquence originale b est immédiate par inversion de 4.21.

Ce schéma a été utilisé pour insérer des messages de 32 bits avec une très bonne robustesse. Pour des raisons qui seront éclaircies au chapitre 6, la décomposition spectrale limite à 7000 sommets la taille des maillages que cette méthode peut traiter (en utilisant [68] pour le calcul des fonctions de base).

4.3.2 Espace d'analyse multirésolution des maillages

Il est possible de définir une analyse multirésolution des maillages [37]. Des schémas de tatouage ont été développés pour utiliser ces bases d'ondelettes [62].

Le message à cacher passe dans une fonction de hachage dont le résultat initialise un générateur pseudo-aléatoire. C'est cette séquence qui sera utilisée comme information utile. Elle est ensuite classiquement ajoutée de manière additive comme un bruit sur le signal décomposé sur les fonctions de base. On commence par construire un vecteur de tatouage w de taille L , utilisé pour perturber les fonctions de base B . Sous forme matricielle, l'insertion s'écrit :

$$\mathcal{V}' = \mathcal{V} + B \times w$$

où B est la matrice des fonctions de base de la transformée en ondelettes. Cette méthode nécessite un alignement préalable du maillage à contrôler sur le maillage original avant de pouvoir relire la marque, de manière non aveugle. Lorsque l'objet a été coupé, il est encore possible de récupérer une synchronisation par corrélation géométrique locale entre les deux objets. Plus généralement, il faut que les deux maillages à contrôler disposent de la même connexité. Si besoin est, on aura recours à un remaillage de l'objet. La marque $\hat{w}$ est estimée en résolvant le système de moindres carrés :

$$B \times \hat{w} = \mathcal{V}_{\mathcal{R}} - \mathcal{V},$$

où $\mathcal{V}_{\mathcal{R}}$ est la géométrie du maillage recalé sur l'original. Afin de tester la présence de la marque, on effectue une corrélation ρ de la marque estimée avec la marque originale :

$$\rho = \frac{\sum_{i \in [1, L]} (\hat{w}_i - \bar{\hat{w}})(w_i - \bar{w})}{\sqrt{\sum_{i \in [1, L]} (\hat{w}_i - \bar{\hat{w}})^2 \times \sum_{i \in [1, L]} (w_i - \bar{w})^2}},$$

où $\bar{x}$ représente la moyenne des éléments du vecteur x . La corrélation ρ est ensuite soumise à un test statistique pour garantir la présence de la marque. Les auteurs rapportent de très bons résultats en robustesse, principalement dus à un réalignement géométrique de bonne qualité sur l'original.

4.3.3 Espace des NURBS

Dans les domaines de la conception assistée par ordinateur, on préfère souvent manipuler des surfaces (ou leurs paramètres) que des maillages. On utilise généralement des NURBS (Non-Uniform Rational B-Splines) pour cela. Dans [57], on trouve plusieurs méthodes de tatouage dédiées à cet espace particulier de représentation des surfaces. Les méthodes varient en fonction des contraintes sur la forme et de celles pesant sur la discrétion de l'insertion. En effet, les surfaces fonctionnelles de l'objet ne tolèrent aucune altération, par contre on pourra jouer sur la forme dans le cas de surfaces libres. En outre, le tatouage peut induire une paramétrisation plus fine de l'objet, nécessitant l'insertion de points de contrôle qui peuvent paraître suspects à un éventuel attaquant. Nous montrons sur la Fig. 4.6 un exemple de spline reparamétrisée en rajoutant des points de contrôle et en changeant la forme.

On procède par reparamétrisation des splines à l'aide d'une fonction rationnelle dont les termes du rapport sont linéaires. Le choix d'une telle fonction permet de garder le même

nombre de points de contrôle et n'alourdit pas l'objet par de l'information due au seul tatouage. Une telle reparamétrisation laisse un seul degré de liberté pour encoder l'information utile (sous contrainte de respecter les limites originales de l'espace paramétrique). On insère quelques bits dans la variation de ce seul paramètre. Cette méthode nécessite donc elle aussi le modèle original pour retrouver la marque. Afin de maximiser la capacité, on répète cette procédure sur toutes les surfaces élémentaires qui sont agencées dans un arbre.

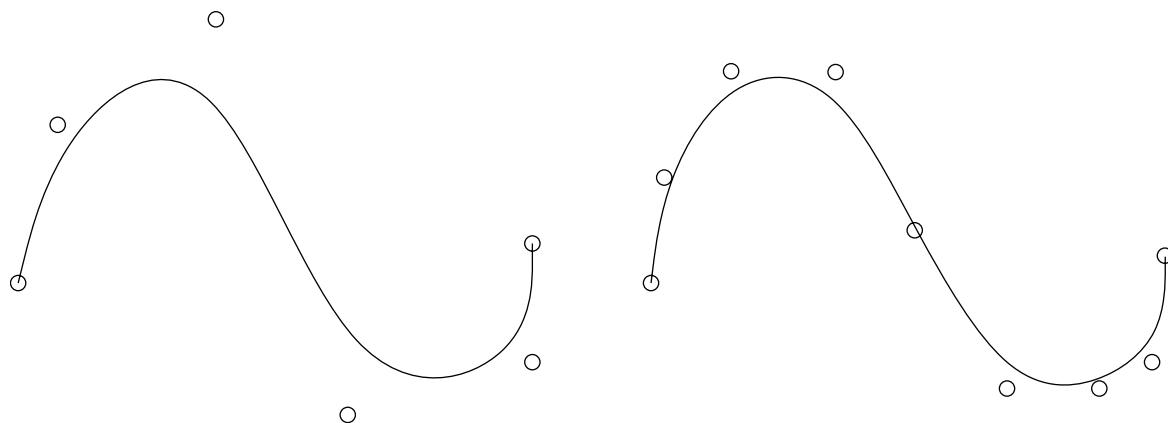


FIG. 4.6 – Tatouage d'une spline en rajoutant des points de contrôle. La forme est supposée libre et a été modifiée.

Une approche étendue permet de doubler la capacité. En effet, on peut étendre la reparamétrisation aux deux splines dont le produit tensoriel forme la surface de l'objet. L'insertion par reparamétrisation concerne alors les deux splines élémentaires, doublant ainsi la capacité.

4.4 Conclusion

La classification que nous avons établie introduit de fait une démarcation par la robustesse entre les différentes méthodes. En particulier, les schémas reposant sur la modification d'un invariant sont plutôt fragiles, et ne tolèrent pas de modification de la connexité, excepté dans le cas d'un arrangement indexé. Ils sont également assez sensibles au bruit. De ce fait, ils paraissent indiqués dans des applications de type authentification. Le chapitre 5 sera pour nous l'occasion de détailler notre approche du tatouage par invariant géométrique. Nous verrons que nous garantirons un certain nombre de spécifications usuelles en tatouage, qui ne sont pas souvent mentionnées ensemble dans la littérature. Nous serons en mesure de donner une probabilité de fausse alarme (ce que seul [62] propose), parce que nous montrerons comment accéder à tous les sites pouvant porter la marque (comme dans [82, 7, 57]), en mettant en place un arrangement indexé (ce qu'utilisent [56] et [82]). Enfin, nous estimerons nous aussi la taille de l'espace des clefs effectivement utilisables (ce que seul [82] garantit). Nous estimerons encore la taille de la plus petite surface protégée, ce qu'à notre connaissance aucune méthode ne mentionne. Enfin, nous expliquerons comment appliquer notre méthode à la stéganographie.

Concernant les autres représentations de l'information 3D des maillages (nous laissons de côté les NURBS), il faut noter qu'il s'agit de représentations autorisant la transmission progressive de la géométrie. Nous détaillerons au chapitre 6 le contexte dans lequel nous avons abordé le tatouage dans l'espace spectral de la géométrie. Nous avons choisi cette représentation car elle permet des analogies fortes avec la transformée de Fourier sur des surfaces régulièrement échantillonnées. Nous montrerons comment modifier l'algorithme de décomposition spectrale pour améliorer la qualité de la transmission progressive de la géométrie. Ce sera l'occasion de développer un codeur de source spectral pour la géométrie. Ce codeur nous permettra de tester notre méthode de tatouage spectral face au codage de source.