

Chapitre 3

Traitements usuels sur les maillages

On souhaite introduire des tatouages robustes aux manipulations de l'objet. Dans cette partie, nous donnons donc une liste de traitements pouvant être appliqués à un objet 3D. Nous distinguons les traitements portant uniquement sur les coordonnées des points, de ceux modifiant également l'information de connexité. Parmi les manipulations n'agissant que sur la géométrie, nous commençons par détailler plus précisément différentes techniques de codage de source.

3.1 Codage de source

Le codage de source [16, 19, 33, 32, 64] est une opération qui peut éventuellement faire varier les coordonnées des points des objets tridimensionnels (compression avec perte). Les différentes familles de codage de source se distinguent par la manière de représenter la connexité. On les regroupe généralement en trois familles :

- celles pour laquelle aucun codage de la connexité n'a lieu (méthodes spectrales),
- celles codant le graphe de la connexité (par exemple la technique préconisée par MPEG-4 [72, 34]),
- celles reposant sur un codage implicite de la connexité (aussi appelées méthodes fondées sur la valence [73, 4, 45, 76]).

Sauf pour les méthodes spectrales, les méthodes de codage de source n'utilisent pas de décorrélation de l'information géométrique : on utilise des schémas de prédiction/correction géométrique. Les méthodes spectrales seront détaillées au chapitre 6.

3.1.1 Schémas de prédiction/correction géométrique

Les points 3D à transmettre sont préalablement ordonnés. On utilise généralement l'une ou l'autre des deux méthodes suivantes pour prédire la position d'un nouveau point à partir de ses antécédents. La première méthode se base sur les deux derniers antécédents (prédiction suivant la règle du triangle) et la seconde sur les trois derniers (prédiction suivant la règle du parallélogramme). La règle de prédiction fondée sur les triangles mesure le vecteur 3D entre le nouveau point à insérer et le milieu du segment formé par ses deux antécédents. La règle de prédiction fondée sur les parallélogrammes mesure le vecteur 3D d'erreur entre la

position du nouveau point à insérer et le point formant un parallélogramme avec les trois derniers antécédents.

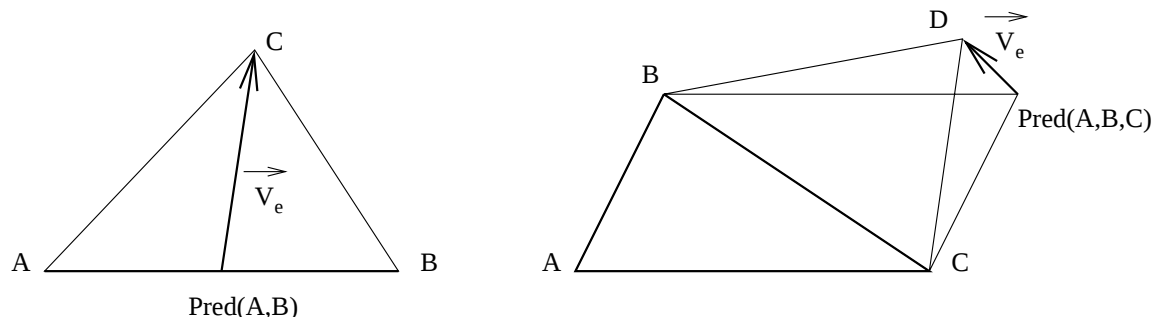


FIG. 3.1 – Schémas de prédiction à partir d'un triangle (à gauche) ou d'un parallélogramme (à droite) : le vecteur d'erreur est mesuré par rapport à la position prédite. En pratique, la prédiction fondée sur la règle du parallélogramme se montre plus efficace : la norme V_e des vecteurs d'erreur est en moyenne plus petite.

Intuitivement, la règle du parallélogramme fonctionne mieux que celle du triangle car la position géométrique de points voisins est généralement corrélée. En pratique la règle du parallélogramme est en effet meilleure que celle du triangle : les vecteurs d'erreur produits sont plus petits en moyenne.

3.1.2 Méthodes fondées sur la valence

Nous présentons deux exemples de méthodes fondées sur la valence : la première propose un codage statique du maillage, la seconde met en œuvre un codage multirésolution. Dans les deux approches, l'idée est que seul l'envoi des valences dans un certain ordre (éventuellement avec quelques codes signalant une exception dans le traitement) suffit à reconstituer la connexité. L'important est de traverser le maillage de manière causale et déterministe, en ne se servant que de l'information de valence.

Méthode de Touma et Gotsman

La première méthode fondée sur la valence [73] propose un encodage statique de la connexité. Elle est valable pour les variétés orientables uniquement. En effet, cette propriété permet d'établir un *ordre local* sur le voisinage d'un sommet, en tournant par exemple dans le sens horaire. Une conséquence de cette propriété est que si un cycle de sommets est retiré du maillage, on crée deux autres maillages, l'un étant à l'intérieur du cycle, et l'autre à l'extérieur. Dans la suite, un tel cycle de sommets sera appelé une liste active. L'idée revient à faire croître cette liste jusqu'à ce que tous les sommets soient parcourus (i.e : soient à l'intérieur du cycle formé par la liste active). Un point à l'intérieur de la liste active est dit conquis. Un point à l'extérieur est dit non conquis. Lorsqu'un point est conquis, la totalité des arêtes adjacentes sont également conquises. Ainsi, un point de la liste active contient des arêtes conquises et non conquises, et est relié par une arête à chacun de ses voisins dans la liste active.

L'encodage de la connexité débute par la sélection d'un triangle dont les trois arêtes forment la liste active initiale. Chaque point de la liste active est passé en revue dans le sens horaire. L'algorithme étend alors la liste active en sélectionnant, dans le sens anti-horaire, les arêtes encore libres autour du point en cours de traitement. Chacune de ces arêtes se voit associé un code *add*, suivi de la valence du sommet non conquis de l'arête en cours de conquête. Une fois toutes les arêtes libres du point conquises, on le place à l'intérieur de la liste active (il est alors conquis). Si la liste active vient à s'intersecter au cours de son expansion, on la découpe alors en deux listes actives, qui seront parcourues successivement. Une telle intersection de la liste active avec elle-même génère un code *split*. L'encodage se termine lorsque toutes les arêtes du maillage sont conquises.

La procédure de décodage est alors triviale : chaque code *add* rencontré provoque l'insertion d'un nouveau point. Ce nouveau point forme alors un nouveau triangle en le connectant avec le point en cours de traitement dans la liste active et à son successeur dans cette liste. Le code de valence permet de déduire le nombre d'arêtes libres pour chaque nouveau sommet dans la liste active.

Il est à noter dès à présent que nous utiliserons une idée semblable dans nos algorithmes de tatouage (au chapitre 5), en rajoutant une contrainte supplémentaire.

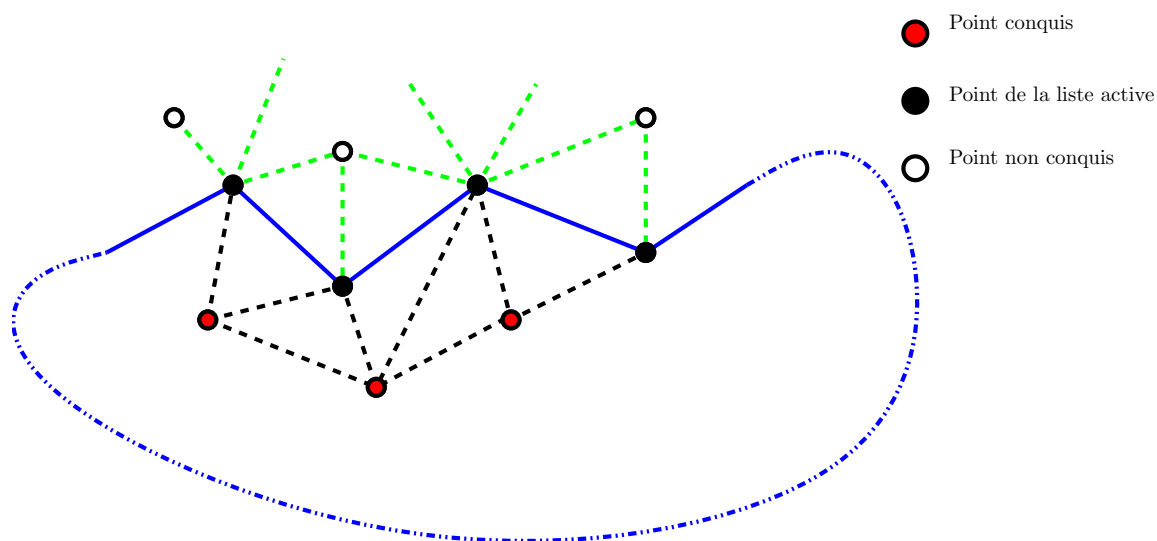


FIG. 3.2 – Propagation sur le maillage grâce à l'information de valence dans la méthode de Touma et Gotsman : il suffit de coder les valences des points que l'on conquiert pour propager la liste active.

Le codage de l'information de géométrie aurait pu être effectué à partir d'une prédiction par la règle du triangle, mais les auteurs ont opté pour une prédiction à base de parallélogramme. Le triangle considéré pour la prédiction est celui formé par le point de la liste active en cours de traitement, son prédécesseur, et le point conquis dont ils sont les voisins.

Méthode d'Alliez et Desbruns

Cette méthode [4] est une extension multirésolution⁷⁵ de l'idée précédente, elle autorise la transmission progressive de l'information géométrique maillée. La différence essentielle est qu'ici, lors du décodage, la rencontre d'un code *add* ne génère pas un triangle mais reconstitue tout le voisinage du nouveau point en cours de décodage. Afin de générer cette liste de points, on commence par paver le maillage avec deux types de cellules : les cellules régulières (constituées d'un point, le centre de la cellule, et de son voisinage) et les cellules dégénérées (constituées d'un seul triangle, n'ayant pas de centre). Chaque cellule dispose également de ports d'entrée/sorties. Un port est constitué d'une arête de la frontière de la cellule, et d'un sommet. Il y a autant de ports que d'arêtes sur la frontière de la cellule. Arbitrairement on impose un seul port d'entrée dans la cellule. Le port d'entrée a son sommet appartenant à la cellule. Tous les autres ports, de sortie de la cellule, ont un sommet hors de la cellule. Des cellules régulières et dégénérées suffisent à paver le maillage.

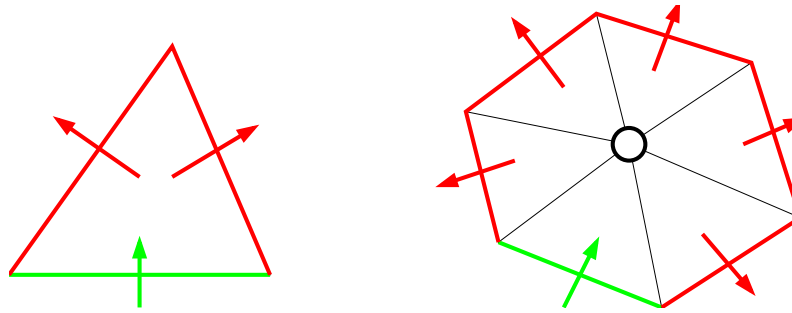


FIG. 3.3 – Les deux types de cellules dans la méthode d'Alliez et Desbrun : à gauche une cellule dégénérée, à droite une cellule régulière. Son centre a ici une valence de 6. Nous avons représenté les ports d'entrée en vert, et les ports de sorties en rouge.

On commence à parcourir le maillage en sélectionnant un port arbitrairement (i.e : on choisit une arête et un des points avec lesquels elle forme un triangle). Ce port est considéré comme un port d'entrée dans une cellule. On sélectionne ainsi implicitement la première cellule, dont on déduit alors tous les ports de sortie. Les sommets du premier port d'entrée sont marqués comme conquis, et le port est placé dans une file. On commence à itérer en extrayant le port de la file, et on examine l'état de son sommet :

- Si le sommet du port est marqué conquis ou à supprimer : il n'y a rien à faire, car la cellule dont ce port est celui d'entrée a déjà été visitée. On écarte ce port, et on sélectionne le suivant dans la file.
- Si le sommet du port est encore libre et si sa valence est inférieure ou égale à d_{max} (valence maximale choisie à 6) : la cellule implicitement désignée par le port d'entrée est régulière. Son centre sera marqué à *supprimer*, supprimé, puis le polygone de son ancien voisinage sera retiangulé de manière implicite. Les sommets de l'ancien voisinage sont marqués *conquis*. On génère alors un symbole *add* suivi de la valence du point dont on

⁷⁵Pour un maillage, le passage d'une résolution à une autre, plus grossière, est caractérisé par la suppression d'un (resp. d'une) ou plusieurs sommets (resp. arêtes). La méthode que nous présentons ici supprime, pour passer à une résolution plus grossière, un ensemble de sommets dont les voisinages ne se recouvrent pas.

vient d'effectuer la suppression. On collecte alors tous les ports de sortie (dans le sens horaire ou anti-horaire), et on les place dans la file. On sélectionne ensuite le prochain port dans la file, et on traitera ainsi la prochaine cellule dont il est le port d'entrée.

- *Sinon (le sommet du port a une valence supérieure à d_{max} ou bien son sommet est marqué conquis)* : nous sommes en présence d'une cellule dégénérée. On déclare *conquis* le sommet du port d'entrée de la cellule dégénérée, et on place ses deux ports de sortie dans la file. Lors de la rencontre d'une cellule dégénérée, on envoie un code spécial signalant la dégénérescence.

Dans l'algorithme précédent, d_{max} est fixé à 6. Nous discuterons cette valeur au fil de la description de l'algorithme. Nous sommes en mesure de paver de manière déterministe le maillage avec des cellules régulières (en majorité), et des cellules dégénérées (plus rares). Pour passer à une résolution plus grossière, on viendra supprimer les centres des cellules régulières et retriangler leur intérieur de manière implicite.

Chaque point conquis se voit attribuer une étiquette positive ou négative s'il faut localement maximiser ou minimiser sa valence, respectivement. On cherche à tendre vers une valence moyenne de 6 sur le maillage décimé (dans un maillage chaque sommet a en moyenne 6 voisins). La Fig. 3.4 expose la manière dont on propage les étiquettes positives et négatives sur la périphérie de l'ancienne cellule régulière. En particulier, chaque sommet marqué positivement (resp. négativement) tendra vers une valence élevée (faible) après la retriangulation. On détermine ainsi la retriangulation de la cellule régulière en fonction de la valence du centre de la cellule et des étiquettes positives et négatives du port d'entrée, puis on propage ces étiquettes.

Lors de la conquête d'une cellule régulière, il se peut qu'un certain nombre des sommets de sa frontière soient déjà marqués avec des étiquettes positives et/ou négatives n'étant compatibles avec aucune configuration canonique. Dans ce cas, on ignore l'anomalie, et on retriangle tout de même suivant la Fig. 3.4. Les étiquettes des sommets conquis ne sont pas modifiées.

Sur la Fig. 3.4 on donne les 16 configurations canoniques pour le remaillage des cellules régulières. En effet, ayant fixé $d_{max} = 6$, il ne reste que 4 possibilités pour la valeur de la valence du centre de la cellule : 3, 4, 5, et 6. Ensuite, le jeu des étiquettes positives et négatives sur le port d'entrée impose lui aussi 4 configurations par type de valence. Ce jeu d'étiquettes devant être implicite, cela implique de faire figurer la convention avec laquelle on le propage en passant d'une résolution à l'autre : après suppression du centre, les étiquettes de la cellule sont entièrement déterminées par la Fig. 3.4.

On voit au passage qu'en faisant ce choix pour la distribution des étiquettes positives et négatives lors de la retriangulation, on tend à faire en sorte qu'un sommet sur deux de la frontière de la cellule régulière ait une valence plus faible. On tend donc à maximiser le nombre de sommets de valence 3, ce qui perturbe la distribution naturelle des valences, centrée autour de 6.

Pour pallier ce problème, on effectue dans la foulée une autre étape de décimation guidée par la valence. Cette fois, on fixera d_{max} à 3 afin de corriger l'effet de décalage des codes de valence vers 3 lors de l'étape précédente. Une niveau de résolution comporte en réalité l'effet de deux conquêtes : la première portant sur les cellules régulières dont le centre a une valence d'au plus 6, et la seconde, de nettoyage, porte uniquement sur les cellules régulières dont le centre a une valence de 3. La décimation de nettoyage impose également de modifier

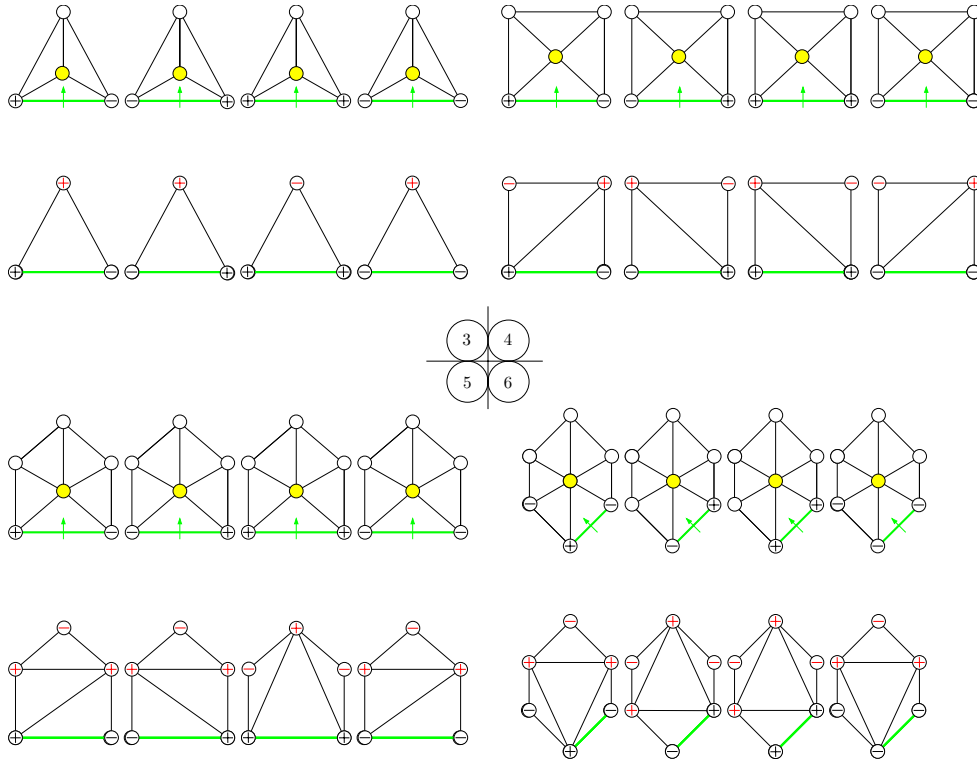


FIG. 3.4 – Triangulations canoniques dans la méthode d’Alliez et Desbrun. Chaque valence possible (3, 4, 5, 6) pour le centre de la cellule régulière engendre quatre jeu d’étiquettes possibles sur le port d’entrée. A chaque fois, il faut définir la retriangulation (du haut vers le bas pour chaque valence) et la propagation des étiquettes.

légèrement les ports de sortie des cellules régulières (dont le centre est à valence 3) : on prendra les quatre ports de sortie dont les arêtes appartiennent aux deux faces adjacentes à la frontière de la cellule régulière.

Pour le décodeur, qui reconstruit les cellules une par une au rythme de l’arrivée des symboles de valence ou de dégénérescence, les étiquettes positives/négatives fournissent le moyen d’une découverte déterministe de la cellule. Les huit cas sont listés sur la Fig. 3.5. Lorsqu’un code de valence $v \geq 3$ arrive, on sait que l’on doit créer $v - 2$ nouvelles facettes. Les huit cas synthétisent comment les déterminer. On reconstitue ainsi de manière déterministe le jeu d’étiquettes pour tous les points de la cellule, afin de pouvoir progresser dans le décodage. On arrive alors à décoder la frontière externe de la cellule uniquement avec l’information de valence et le jeu d’étiquettes du port d’entrée. C’est ainsi que le codeur et le décodeur sont correctement synchronisés : grâce à la propagation implicite des étiquettes tant au codage qu’au décodage. Cette approche fondée sur un parcours implicite a pu encore être généralisée aux maillages non nécessairement triangulés, et aux données volumétriques hexaédriques.

On arrive ainsi à représenter hiérarchiquement et de manière déterministe un maillage, en utilisant uniquement l’information de valence. L’optimalité de ce codage a été discutée dans [45], il tend à se rapprocher d’un optimum énoncé par Tutte [75] dans le cas des maillages planaires.

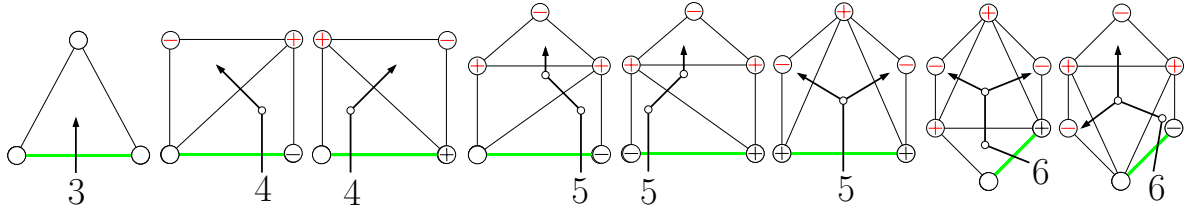


FIG. 3.5 – Découverte déterministe de la cellule régulière par l'information de valence dans la méthode d'Alliez et Desbrun. Etant donné le port d'entrée (avec ses étiquettes) dans la cellule régulière à créer, le code de la valence $v \geq 3$ suffit à déterminer comment créer les $v - 2$ nouveaux triangles. Les codes laissés en blanc signifient que le signe de l'étiquette est indifférent.

Prédire correctement la position du centre des cellules régulières est l'objet de la compression de l'information géométrique. Pour cela, les auteurs se placent dans une base de Frenet locale à la cellule et prédisent les erreurs radiale et tangentielle sur la position du centre de la cellule.

Méthode de Valette (ondelettes surfaciques)

La méthode de Valette [76] est une généralisation de la méthode de Lounsbery [52] pour les maillages de subdivision. Un maillage est dit de subdivision lorsqu'il est le résultat d'un processus itéré de subdivisions régulières partout ou par morceaux. Nous donnons sur la Fig. 3.6 un exemple de prédiction par subdivision régulière pour deux triangles. Cette méthode est limitée car elle impose une contrainte forte sur la connexité du maillage. La subdivision

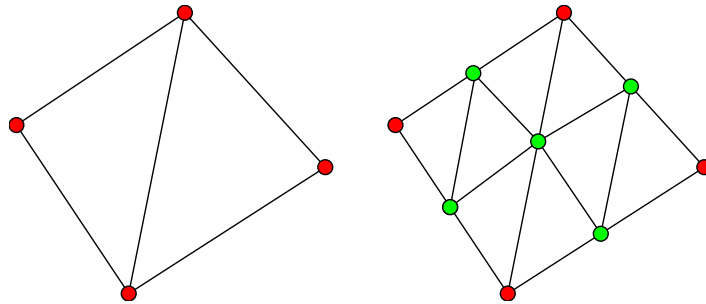


FIG. 3.6 – Exemple de subdivision régulière d'un triangle en quatre autres triangles (subdivision de Loop [50]).

consiste à prédire la position de points, pour ensuite ne stocker que l'erreur commise. On place alors du mieux qu'on peut les prédictions des nouveaux points. Dans la Fig. 3.6, ils sont placés au milieu des arêtes des triangles à subdiviser. On peut alors écrire naturellement le passage d'un niveau d'analyse j de maillage $\mathcal{S}^j(\mathcal{V}^j, \mathcal{E}^j)$ au suivant $(j + 1)$ par un système de deux filtres d'analyse $\mathcal{A}^j$ et $\mathcal{B}^j$, et la reconstruction s'effectue à l'aide des filtres $\mathcal{Q}^j$ et $\mathcal{R}^j$. Dans ce cas précis d'un maillage de subdivision, il n'y a pas d'innovation de l'information de

connexité : la subdivision est unique. On passe implicitement de la connexité de $\mathcal{E}^j$ à celle de $\mathcal{E}^{j+1}$. L'analyse s'écrit donc :

$$\mathcal{V}^{j+1} = \mathcal{A}^j \times \mathcal{V}^j, \quad (3.1)$$

$$\mathcal{D}^{j+1} = \mathcal{B}^j \times \mathcal{V}^j, \quad (3.2)$$

où les $\mathcal{D}^j$ sont les coefficients d'ondelettes permettant de reconstruire $\mathcal{V}^{j+1}$ avec $\mathcal{V}^j$ par :

$$\mathcal{V}^{j+1} = \mathcal{Q}^j \times \mathcal{V}^j + \mathcal{R}^j \times \mathcal{D}^j. \quad (3.3)$$

Nous représentons les vecteurs d'erreur à coder sur la Fig. 3.7.

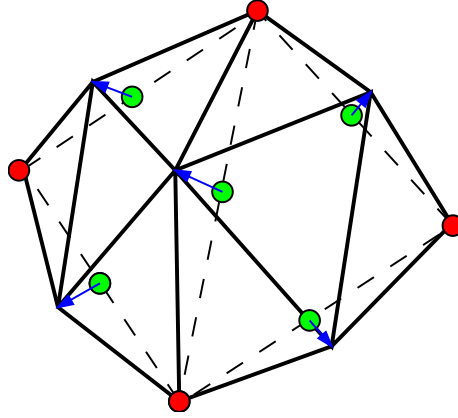


FIG. 3.7 – Vecteurs d'erreur à coder dans l'hypothèse d'une analyse multirésolution par subdivision régulière de Loop.

On peut généraliser cette technique restreinte aux maillages de subdivision grâce à une subdivision plus riche, irrégulière, permettant d'appréhender n'importe quelle connexité. Cette subdivision irrégulière recense 15 configurations canoniques (et dont nous donnons quelques exemples sur la Fig. 3.8), et elle autorise la subdivision d'un triangle en 4, 3, 2 triangles, ou le laisser inchangé.

A l'encodage, il faut donc trouver un moyen d'apparier les facettes entre elles du mieux possible afin de ne jamais bloquer l'algorithme. Les auteurs proposent une manière d'y parvenir pouvant éventuellement faire appel à une convention faisant d'une arête une facette spéciale, dégénérée. On donne ci-dessous (voir Fig. 3.9) un exemple de fusion des facettes opérée par l'algorithme.

Une fois les facettes groupées, on envoie un code décrivant la configuration de la subdivision irrégulière. Afin de procéder à une réelle analyse par ondelettes, les auteurs complètent cette technique par un lifting, qui permet d'obtenir à la résolution j la meilleure approximation $\mathcal{V}^j$ de $\mathcal{V}^{j+1}$ au sens des moindres carrés [60].

3.1.3 Méthode de Taubin (MPEG-4)

Ce codage est celui employé dans la compression du format binaire VRML et dans MPEG-4 3D Mesh Coding [72]. Il ne s'applique qu'aux maillages de connexité uniquement triangulaire. Le principe en est le suivant : le maillage est, si nécessaire, découpé de telle manière

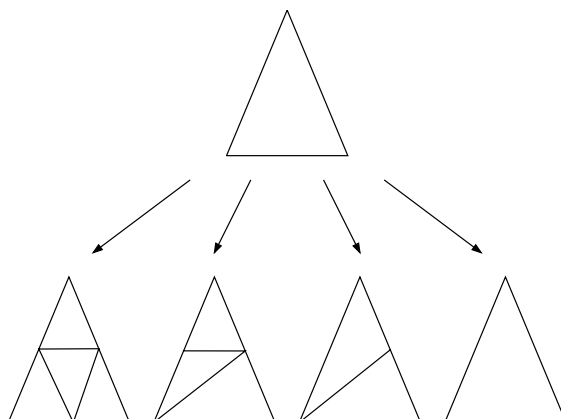


FIG. 3.8 – Quelques cas possibles parmi ceux d’une subdivision irrégulière d’un triangle en un nombre variable d’autres triangles.

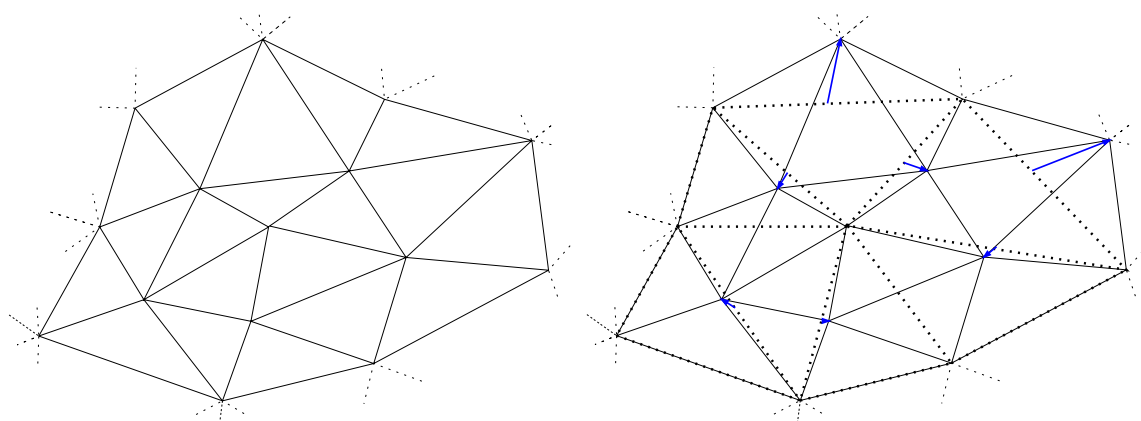


FIG. 3.9 – Fusion des facettes dans la méthode de Valette : une heuristique se charge de regrouper les facettes avant décimation.

que chaque arête appartienne au plus à deux facettes [34]. Ensuite, chaque composante est littéralement pelée en longs rubans de triangles de manière à obtenir un arbre couvrant des faces. A ce premier arbre est adjoint un arbre de recouvrement des sommets. Ainsi, le graphe de connexité “pelé” peut-il se projeter sur un plan 2D. Le graphe alors projeté sur le plan se réduit à un polygone simple qu’il est possible de coder de manière extrêmement compacte avec un alphabet de quatre symboles. Les deux arbres de recouvrement contiennent toute l’information nécessaire à la reconstruction de la connexité originale.

La compression de la géométrie a alors lieu sur les longues bandes de triangles par prédiction et correction. Un codage correct de la géométrie nécessite un minimum de 10 bits par échantillon pour obtenir un effet visuel satisfaisant avec une prédiction fondée sur les parallélogrammes. Cette méthode de compression est la seule faisant l’objet d’une normalisation (VRML 2.0 Binaire et MPEG-4 3D Mesh Coding).

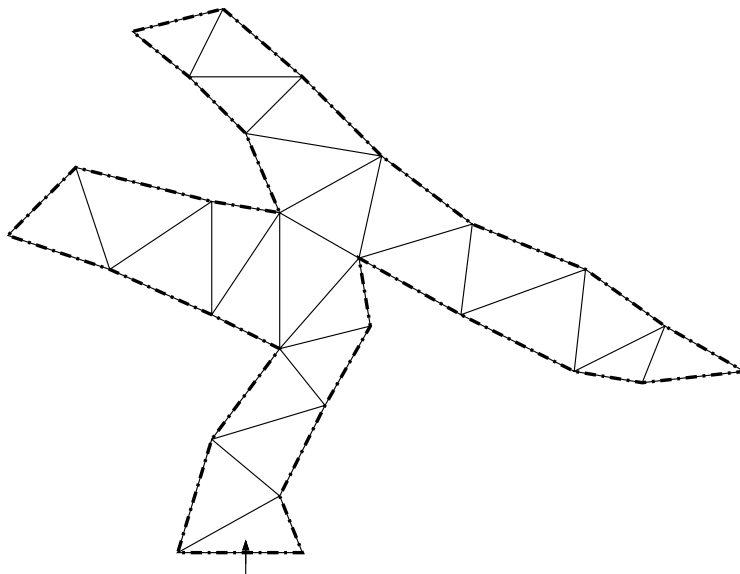


FIG. 3.10 – Encodage de la connexité dans MPEG-4 : le maillage est pelé, puis projeté sur le plan afin d’obtenir un polygone simple. On part alors d’un triangle et d’une direction d’entrée pour commencer à décrire l’arbre de connexité des triangles avec l’alphabet de quatre symboles nécessaires pour fixer les quatre situations possibles relativement au(x) triangle(s) suivant(s) : G (resp. D) si son unique fils est à gauche (resp. droite), 2 s’il a deux fils et F s’il est une feuille. Muni d’une orientation ou d’un sens de parcours, l’arbre en parcouru en profondeur d’abord. Ici le code de connexité est : GDGDGG22DGDGDGFDGDGFDGDGF.

3.2 Traitements affectant la géométrie uniquement

Si le codage de source demeure le principal traitement subi par les signaux, il n’en reste pas moins que le tatouage peut éventuellement souffrir d’autres opérations. Ce sont typiquement celles proposées dans la plupart des logiciels de modélisation 3D. Nous en donnons à présent une liste non exhaustive. Dans chaque cas, nous discutons l’impact de la manipulation considérée sur le schéma de tatouage.

3.2.1 Ajout de bruit

L’ajout de bruit sur les coordonnées des points est souvent le résultat de la transmission sur des canaux médiocres. C’est aussi l’une des attaques classiques de ceux qui cassent les tatouages ; c’est enfin une façon de modéliser la compression géométrique.

3.2.2 Lissage

Lors de l’acquisition des données 3D par un scanner, il arrive couramment que la surface présente un aspect rugueux. Afin d’éliminer cet artefact des techniques de lissage ont été proposées comme celle exposée dans [71]. Taubin reprend la définition du filtrage gaussien [49, 81] et propose un opérateur laplacien défini sur les surfaces maillées. A chaque point p_i

on ajoute son laplacien pondéré par λ :

$$p'_i = p_i + \lambda \Delta p_i, \quad (3.4)$$

avec $0 \leq \lambda \leq 1$. Le laplacien de Taubin est défini en posant :

$$\Delta p_i = \sum_{j \in i^*} w_{ij} (p_j - p_i). \quad (3.5)$$

On impose encore :

$$\sum_{j \in i^*} w_{ij} = 1. \quad (3.6)$$

On a donc :

$$p'_i = (1 - \lambda)p_i + \lambda \sum_{j \in i^*} w_{ij} p_j, \quad (3.7)$$

On choisit alors des poids $w_{i,j}$ fonctions des arêtes liant p_i à ses voisins :

$$w_{ij} = \frac{\phi(p_i, p_j)}{\sum_{k \in i^*} \phi(p_i, p_k)}. \quad (3.8)$$

où d_i est la valence du sommet $p - i$. Les fonctions ϕ ne dépendant que des arêtes, on peut proposer plusieurs choix :

$$\phi(i, j) = \frac{1}{d_i}, \quad (3.9)$$

$$\phi(i, j) = \|p_i - p_j\|^\alpha. \quad (3.10)$$

Un autre choix consiste à prendre pour $\phi(i, j)$ le rapport de la somme des aires des deux faces adjacentes à l'arête $\{i, j\}$ sur deux fois l'aire du voisinage de p_i . Nous illustrons un tel filtrage gaussien, avec $\phi(i, j) = \frac{1}{d_i}$ et $\lambda = 0.2$, sur la Fig. 3.11. Nous avons effectué 4000 itérations. Nous reviendrons sur cet opérateur dans le chapitre consacré à la décomposition spectrale.

3.2.3 Rotation, translation et mise à l'échelle uniforme

Il s'agit ici des manipulations le plus souvent appliquées aux objets 3D. Sous cette appellation on regroupe les transformations affines : la rotation, la mise à l'échelle et la translation, ou encore une combinaison des trois. Afin d'être robuste face à ces transformations, les schémas de tatouage doivent en fait les prendre en compte dès leur conception. Dans la grande majorité des cas, on utilise des primitives d'insertion de la marque naturellement invariantes par transformation affine.

3.2.4 Transformations perspectives

A la différence des précédentes, les transformations perspectives ne sont pas linéaires. Elles sont utilisées pour déterminer l'image d'un objet 3D par un système optique et sont disponibles dans les modeleurs commerciaux. Là encore, on privilégierait un espace de tatouage naturellement invariant. Toutefois comme les transformations perspectives transforment des données 3D en données 2D, on perd généralement le tatouage.

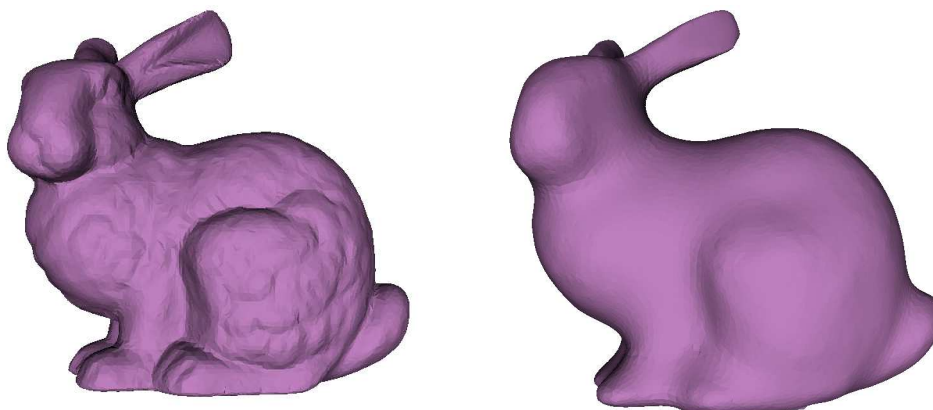


FIG. 3.11 – Exemple de maillage exagérément lissé. A gauche : original (données : Stanford 3D Scanning Repository). A droite : maillage lissé.



FIG. 3.12 – Exemple de transformation perspective. Données : Stanford 3D Scanning Repository.

3.3 Traitements affectant la connexité

Nous présentons à présent une liste non exhaustive de quelques traitements modifiant l'information de connexité. L'information géométrique peut même varier également. Du point de vue du tatouage, cette classe de transformations complique passablement la conception d'un procédé d'insertion robuste de la marque.

3.3.1 Renumerotation des points

Après certains traitements supposant une modification de la connexité (décimation, remaillage), les points sont renumérotés, certains ayant disparu ou leur nombre ayant changé. Certains schémas primitifs de tatouage pouvaient utiliser l'index du point considéré pour inscrire un bit à cacher, ils étaient donc très sensibles à la renumérotation. Toutefois, la

plupart des schémas actuels n'utilisent plus cette information.

3.3.2 Coupe

La coupe, ou section, est une modification de la connexité consistant à extraire telle ou telle partie du maillage. C'est l'équivalent du fenêtrage (*crop*) pour le cas de l'image fixe. La géométrie et la connexité des sommets restants sont inchangées. Les schémas de tatouage se servant de la connexité pour situer l'information cachée peuvent être mis en défaut par ce type de manipulation.

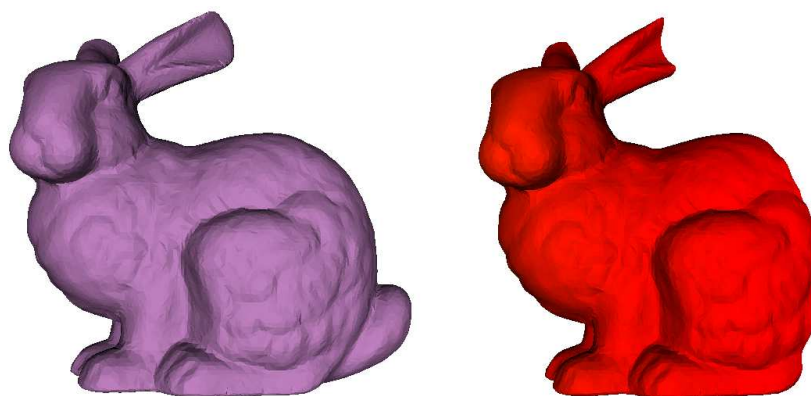


FIG. 3.13 – A gauche : original. A droite : exemple de coupe (par un plan).

3.3.3 Décimation

Les objets 3D correspondent à un type de données complexe. Ils sont donc extrêmement lourds à manipuler. Au début de l'imagerie 3D, se posait le problème d'un affichage rapide de tels objets sur des dispositifs aux capacités limitées. En observant que la suppression de quelques points ou arêtes judicieusement choisis ne dénaturait pas l'aspect visuel de l'objet, la communauté s'est mise à la recherche de moyens de simplifier les maillages, de les décimer. Cet axe de recherche a été particulièrement fécond et de nombreux travaux portent sur ce domaine. On distingue deux sortes d'algorithmes de décimation : ceux procédant par l'effondrement de points, et ceux supprimant des arêtes (voir Fig. 3.14). Tous ces algorithmes ont en commun de proposer une méthode déterministe de remaillage local. La plupart du temps, le choix du point ou de l'arête à effondrer s'opère en fonction de l'évaluation d'une fonction d'énergie locale. On choisit généralement de minimiser une forme quadratique comme dans l'exemple que nous développons ci-dessous.

Dans [31], une technique de décimation appelée QEM (pour Quadric Error Metric) permet de simplifier un maillage à l'aide de l'opérateur d'effondrement d'arête. Chaque itération de l'algorithme classe les arêtes par priorité. La priorité est donnée aux effondrements d'arête minimisant une certaine erreur. L'erreur $E(p'_i)$ commise en remplaçant p_i , appartenant à un plan de normale $\vec{n}$ et de constante d , par le point p'_i , est calculée comme étant la distance

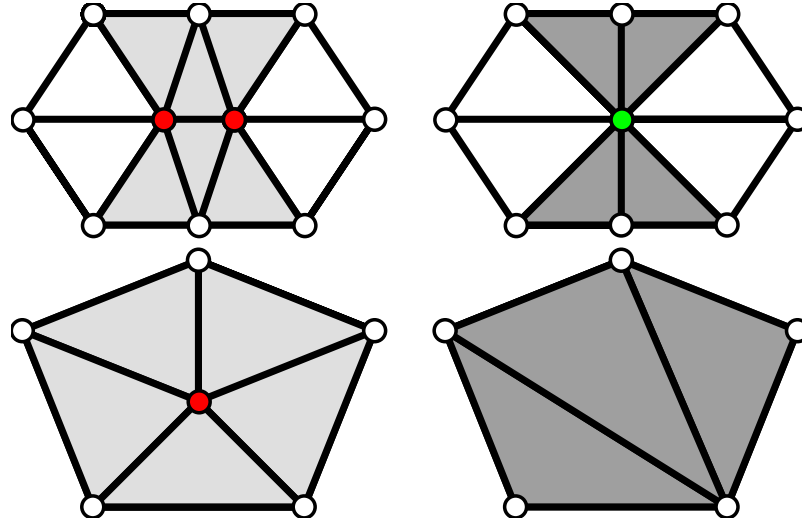


FIG. 3.14 – Principaux opérateurs de décimation (également appelés opérateurs d'Euler) entraînant une modification locale de la connexité. En haut : effondrement d'arête. En bas : effondrement de point.

au carré du point au plan :

$$E(p'_i) = (\vec{n}^T p'_i + d)^2 = p_i'^T (\vec{n} \vec{n}^T) p'_i + 2d \vec{n}^T p'_i + d^2. \quad (3.11)$$

Il s'agit d'une forme quadratique avec un terme linéaire et une constante. On met le tout sous la forme :

$$E(p'_i) = p_i'^T A p'_i + 2b^T p'_i + c. \quad (3.12)$$

Finalement, l'erreur totale commise en un point en remplaçant p_i par p'_i est calculée comme la somme des erreurs commises par rapport aux plans des facettes du voisinage de p_i . L'erreur en un point p_i est minimisée pour p_i qui est le seul à appartenir à tous les plans des facettes de son voisinage, voir Fig. 3.15.

La décimation consiste à trouver un point $\bar{p}$ remplaçant au mieux l'arête $\{p_i, p_j\}$ à effondrer. L'erreur commise est donc fonction en réalité des deux points extrémaux de l'arête. Pour une approximation du calcul de l'erreur commise en remplaçant une arête par un point, les auteurs proposent de faire la somme des erreurs en chaque point de l'arête à effondrer. C'est en ce sens que l'on arrive à classer les arêtes par priorité. Bien entendu, la file de priorité est remise à jour après chaque effondrement.

Minimiser l'erreur commise en remplaçant p_i par p'_i revient à résoudre un système linéaire. La solution est donnée par :

$$p_i'^{min} = -A^{-1}b. \quad (3.13)$$

Et l'erreur minimale ainsi commise vaut :

$$E_{min} = E(p_i'^{min}) = -b^T A^{-1}b + c. \quad (3.14)$$

Du point de vue du tatouage, la modification de l'objet porte à la fois sur les informations de connexité et de géométrie. C'est donc toute la *représentation* de l'objet qui change.

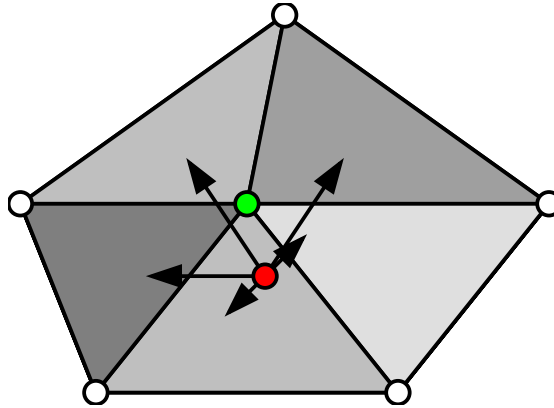


FIG. 3.15 – Calcul de l'erreur quadratique en un point dans QEM : c'est la somme des distances au carré d'un point aux plans des facettes.



FIG. 3.16 – Exemple de simplification uniforme. A gauche : original. A droite : version simplifiée (90% des points effondrés).

Pourtant, la géométrie au sens large ne s'en trouve pas nécessairement bouleversée. Nous illustrons cette méthode de simplification à la Fig. 3.16 où 90% des points ont été supprimés.

3.3.4 Subdivision

Le problème de la subdivision prend encore sa source en visualisation. Une fois le maillage simplifié, on cherche un moyen d'augmenter son nombre de points de manière automatique, sans recourir aux sommets originaux. Une procédure permettant de rajouter des points se révèle en machine plus efficace que le stockage et l'affichage des points réels.

On distingue les procédures de subdivision d'approximation (Loop [50]) et d'interpolation (Butterfly modifié [20, 83]). Nous ne présentons que les subdivisions régulières définies sur des triangles, même s'il en existe pour des quadrilatères (Catmull-Clark pour l'approximation [14] et Kobbelt pour l'interpolation [46]). Les subdivisions d'interpolation sont beaucoup

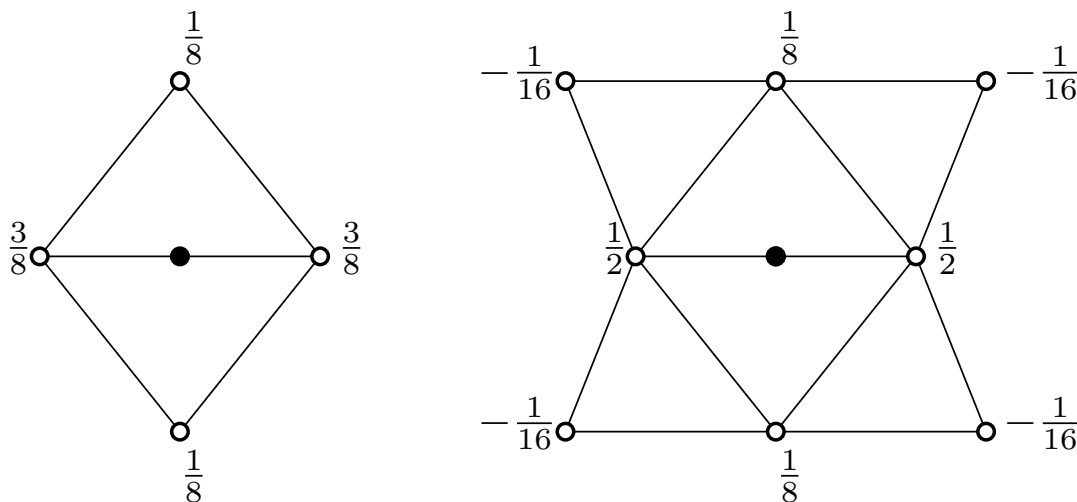


FIG. 3.17 – Voisinages et poids pour le calcul des coordonnées d'un nouveau sommet dans les schémas de subdivision de Loop (à gauche) et Butterfly modifié (à droite) pour l'interpolation. Les sommets à rajouter sont en noir.

plus souples à manier (puisque les sommets originaux appartiendront au maillage final), mais donnent généralement des surfaces de moins bonne qualité que les subdivisions par approximation.

Une fois la procédure de subdivision choisie, on s'intéresse généralement à la surface limite obtenue par un tel procédé. En particulier, on cherche à caractériser la dérivabilité de la surface limite. Elle est par construction de classe C^2 partout sauf sur les lignes de crête et les bords, où elle peut être de classe C^1 [66].

Généralement, une procédure de subdivision rajoute un sommet au milieu d'une arête (voir Fig. 3.6) et indique comment placer correctement ce sommet en fonction de son voisinage. Les coordonnées du nouveau sommet sont calculées en modulant les coordonnées de son voisinage par des poids idoines. Ensuite, le nouveau sommet est connecté à ses nouveaux voisins, et le processus est itéré jusqu'à obtenir le nombre de sommets souhaité. Nous donnons sur la Fig. 3.17 les voisinages et les poids des subdivisions de Loop et du schéma Butterfly modifié.

3.3.5 Remaillage

Il arrive que certains traitements complexes que l'on ait à effectuer sur des maillages requièrent une connexité particulière (typiquement une connexité régulière avec des valences égales à 6, ou une connexité de subdivision [78]).

Un pré-traitement est alors nécessaire : on vient remailler l'objet avec la connexité souhaitée. Un remaillage peut également être utile dans le cas où l'on souhaite répartir judicieusement les points en fonction de la courbure de l'objet [5]. Là non plus, la géométrie au sens large n'est que peu affectée par ces opérations. En revanche, la représentation entière de l'objet est modifiée. On distingue généralement deux approches en remaillage : celle qui modifie le maillage original [74, 38, 47, 12], et celle qui s'en sert pour guider la création d'un

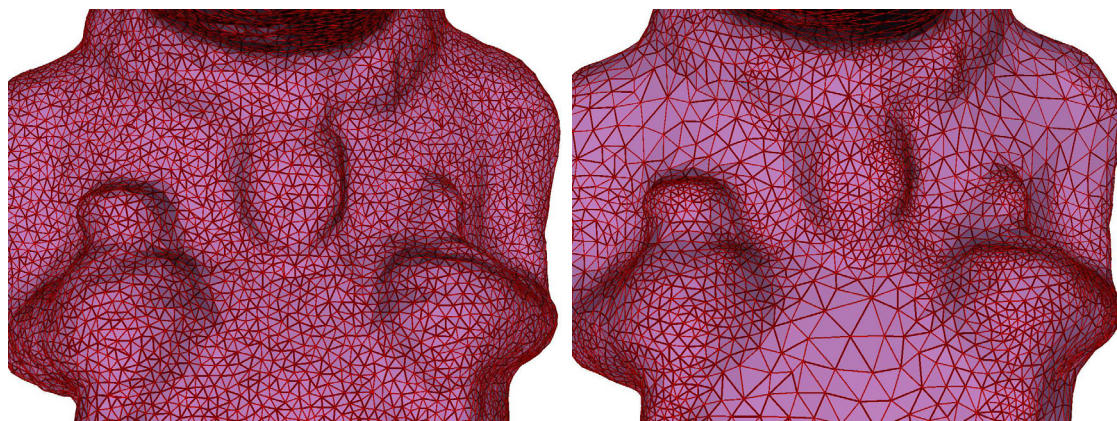


FIG. 3.18 – Exemple de remaillage. A gauche : original, maillage uniforme. A droite : version remaillée, avec une densité adaptée à la courbure. Données : Télécom Paris [36].

maillage neuf [5, 2, 3]. Il s'agit d'une des manipulations les plus sévères qu'un schéma de tatouage ait à affronter.

3.4 Conclusion

Il existe tant de manipulations sur les surfaces maillées que nous avons préféré mentionner les principales, et les séparer en deux catégories suivant qu'elles modifient ou non la connexité de la surface originale. En effet, du point de vue du tatouage, il arrive que l'on se serve de l'information de connexité pour insérer et relire l'information. Nos travaux en tatouage par invariants se placent par exemple dans cette optique. Mais la même surface pouvant être représentée de multiples façons, il se pose le problème de la représentation canonique de cette surface. Une telle représentation n'existe pas en général, et on adopte la représentation la plus adaptée au but à atteindre : en visualisation on cherchera à simplifier le maillage, en simulation on cherchera à uniformiser les éléments de surface, en traitement de la géométrie on cherchera à atteindre une connexité de subdivision.

