



(12) BẢN MÔ TẢ SÁNG CHẾ THUỘC BẢNG ĐỘC QUYỀN SÁNG CHẾ

(19) Cộng hòa xã hội chủ nghĩa Việt Nam (VN)
CỤC SỞ HỮU TRÍ TUỆ

(11)



1-0052227

(51)^{2020.01} H04N 19/00; H04N 19/132; H04N 19/70; H04N 19/176; H04N 19/513; H04N 19/577; H04N 19/105; H04N 19/139 (13) B

(21) 1-2022-02084

(22) 09/10/2020

(86) PCT/US2020/055153 09/10/2020

(87) WO2021/072326 15/04/2021

(30) 62/913,141 09/10/2019 US

(45) 27/10/2025 451

(43) 25/07/2022 412A

(73) BEIJING DAJIA INTERNET INFORMATION TECHNOLOGY CO., LTD. (CN)
Room 101D1-7, 1st Floor, Building 1, No. 6, Shangdi West Road, Haidian District,
Beijing 100085, China

(72) XIU, Xiaoyu (CN); CHEN, Yi-Wen (CN); WANG, Xianglin (US); YU, Bing (CN).

(74) Công ty Luật TNHH T&G (TGVN)

(54) PHƯƠNG PHÁP TÍNH CHỈNH DỮ ĐOÁN VỚI LUỒNG QUANG HỌC, THIẾT
BỊ ĐIỆN TOÁN VÀ PHƯƠNG TIỆN LƯU TRỮ ĐƯỢC ĐỌC ĐƯỢC BỞI MÁY TÍNH
KHÔNG CHUYÊN TIẾP

(21) 1-2022-02084

(57) Sáng chế đề cập đến phương pháp tinh chỉnh dự đoán với luồng quang học (Prediction Refinement with Optical Flow, PROF), thiết bị điện toán và phương tiện lưu trữ đọc được bởi máy tính không chuyên tiếp. Bộ giải mã thu nhận ảnh tham chiếu I được kết hợp với khối video trong tín hiệu video. Bộ giải mã thu nhận các mẫu dự đoán $I(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu I . Bộ giải mã điều khiển các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng việc dịch chuyển sang phải cho các thông số PROF bên trong dựa trên giá trị dịch chuyển bit để đạt được sự chính xác được thiết lập trước. Bộ giải mã thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên quy trình dẫn xuất PROF đang được áp dụng cho khối video dựa trên các mẫu dự đoán $I(i, j)$. Bộ giải mã thu nhận các mẫu dự đoán của khối video dựa trên tổ hợp của các mẫu dự đoán và các giá trị tinh chỉnh dự đoán.

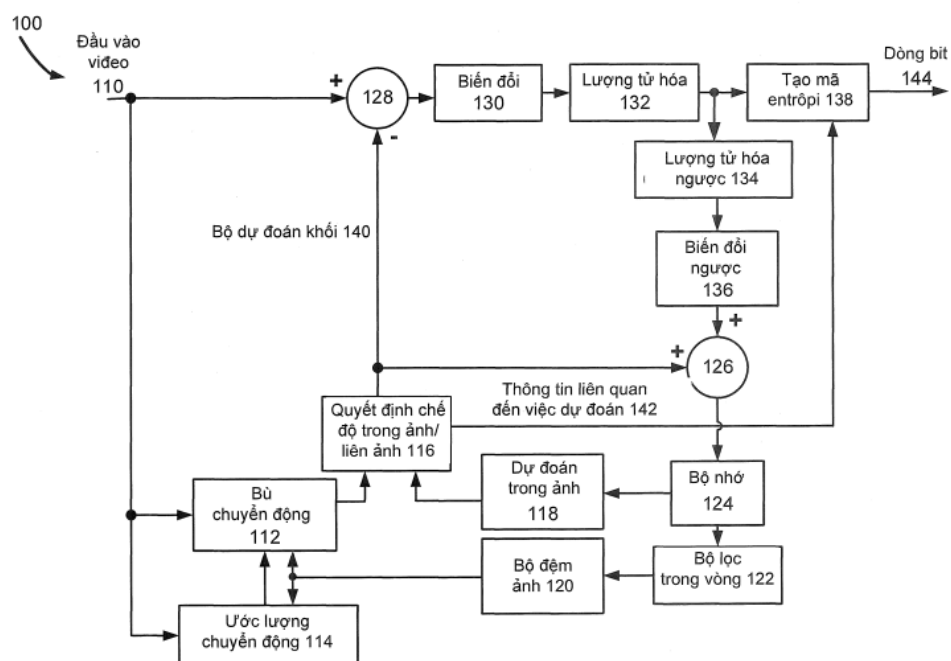


FIG. 1

Lĩnh vực kỹ thuật được đề cập

Sáng chế này liên quan đến việc tạo mã và nén video. Cụ thể hơn là, sáng chế này liên quan đến các phương pháp và thiết bị trên hai công cụ dự đoán liên ảnh mà được nghiên cứu trong tiêu chuẩn tạo mã video vạn năng (Versatile Video Coding, VVC), tức là, sự tinh chỉnh dự đoán với luồng quang học (Prediction Refinement with Optical Flow, PROF) và luồng quang học hai chiều (Bi-Directional Optical Flow, BDOF).

Tình trạng kỹ thuật của sáng chế

Các kỹ thuật tạo mã video khác nhau có thể được sử dụng để nén dữ liệu video. Việc tạo mã video được thực hiện theo một hoặc nhiều tiêu chuẩn mã hóa video. Ví dụ, các tiêu chuẩn tạo mã video gồm việc tạo mã video vạn năng (VVVC), tạo mã mô hình kiểm tra thăm dò hợp tác (Joint Exploration Test Model Coding, JEM), tạo mã video hiệu quả cao (High-Efficiency Video Coding, H.265/HEVC), tạo mã video cải tiến (H.264/Advanced Video Coding, AVC), tạo mã nhóm chuyên gia hình ảnh chuyển động (Moving Picture Experts Group Coding, MPEG), hoặc tương tự. Việc tạo mã video thường tận dụng các phương pháp dự đoán (ví dụ, việc dự đoán liên thông, việc dự đoán nội, hoặc dạng tương tự) tận dụng ưu điểm của sự dư thừa có mặt trong các hình ảnh hoặc các chuỗi video. Mục tiêu quan trọng của các kỹ thuật mã hóa video là để nén dữ liệu video thành dạng sử dụng tỉ lệ bit thấp hơn, trong khi xóa bỏ hoặc tối thiểu hóa các suy giảm chất lượng video.

Bản chất kỹ thuật của sáng chế

Các ví dụ của sáng chế đề xuất các phương pháp và thiết bị để điều khiển độ sâu bit cho luồng quang học hai chiều (BDOF).

Theo khía cạnh thứ nhất, sáng chế đề xuất phương pháp điều khiển độ sâu bit của việc tạo mã tín hiệu video. Phương pháp có thể gồm bộ giải mã thu nhận ảnh tham chiếu I được kết hợp với khối video trong tín hiệu video. Bộ giải mã cũng có thể thu nhận các mẫu dự đoán $I(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu I . i và j biểu diễn tọa độ của một mẫu với khối video. Bộ giải mã còn có thể điều khiển các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng việc dịch chuyển sang phải cho các thông số PROF bên trong dựa trên giá trị dịch chuyển bit để đạt được sự chính xác được thiết lập trước. Các thông số PROF bên trong có thể gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu dự đoán $I(i, j)$. Bộ giải mã cũng có thể thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên quy trình dẫn xuất PROF đang được áp dụng cho khối video dựa trên các mẫu dự đoán $I(i, j)$. Bộ giải mã còn có thể thu nhận các mẫu dự đoán của khối video dựa trên tổ hợp của các mẫu dự đoán và các giá trị tinh chỉnh dự đoán.

Theo khía cạnh thứ hai, sáng chế đề xuất phương pháp điều khiển độ sâu bit của BDOF. Phương pháp có thể gồm bước bộ giải mã thu nhận ảnh tham chiếu thứ nhất $I^{(0)}$ và ảnh tham chiếu thứ hai $I^{(1)}$ được kết hợp với khối video. Ảnh tham chiếu thứ nhất $I^{(0)}$ ở trước ảnh hiện tại và ảnh tham chiếu thứ hai $I^{(1)}$ ở sau ảnh hiện tại trong thứ tự hiển thị. Bộ giải mã cũng có thể thu nhận các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ nhất $I^{(0)}$. i và j có thể biểu diễn tọa độ của một mẫu với ảnh hiện tại. Bộ giải mã cũng có thể thu nhận các mẫu dự đoán thứ hai

$I^{(1)}(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ hai $I^{(1)}$. Bộ giải mã cũng có thể điều khiển các thông số BDOF bên trong của quy trình dẫn xuất BDOF bằng cách áp dụng việc dịch chuyển cho các thông số BDOF bên trong. Các thông số BDOF bên trong bao gồm các giá trị gradient ngang và các giá trị gradient dọc được dẫn xuất dựa trên các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$, các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, các sự chênh lệch mẫu giữa các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ và các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, và các thông số dẫn xuất BDOF trung gian. Các thông số dẫn xuất BDOF trung gian có thể gồm các thông số sGxdI, sGydI, sGx2, sGxGy, và sGy2. sGxdI và sGydI có thể gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị chênh lệch mẫu, và giữa các giá trị gradient dọc và các giá trị chênh lệch mẫu. sGx2 và sGy2 có thể gồm các giá trị tương liên tự động của các giá trị gradient ngang và các giá trị gradient dọc. sGxGy có thể gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị gradient dọc.

Theo khía cạnh thứ ba, sáng chế đề xuất phương pháp của BDOF, PROF, và DMVR. Phương pháp có thể gồm bước bộ giải mã nhận ba cờ điều khiển trong tập hợp thông số chuỗi (Sequence Parameter Set, SPS). Cờ điều khiển thứ nhất chỉ ra liệu BDOF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ hai chỉ ra liệu PROF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ ba chỉ ra liệu DMVR có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Bộ giải mã cũng có thể nhận cờ có mặt thứ nhất trong SPS khi cờ điều khiển thứ nhất là đúng, cờ có mặt thứ hai trong SPS khi cờ điều khiển thứ hai là đúng và cờ có mặt thứ ba trong SPS khi cờ điều khiển thứ ba là đúng. Bộ giải mã còn có thể nhận cờ điều khiển ảnh thứ nhất trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ nhất trong SPS chỉ ra rằng BDOF không được cho phép cho các khối video ở trong ảnh. Bộ giải mã cũng có thể nhận cờ

điều khiển ảnh thứ hai trong phần đầu ảnh của mỗi ảnh khi còn có mặt thứ hai trong SPS chỉ ra rằng PROF không được cho phép cho các khối video ở trong ảnh. Bộ giải mã còn có thể nhận còn điều khiển ảnh thứ ba trong phần đầu ảnh của mỗi ảnh khi còn có mặt thứ ba trong SPS chỉ ra rằng DMVR không được cho phép cho các khối video ở trong ảnh.

Theo khía cạnh thứ tư, sáng chế đề xuất thiết bị điện toán. Thiết bị điện toán có thể gồm một hoặc nhiều bộ xử lý, bộ nhớ đọc được bởi máy tính không chuyển tiếp lưu trữ các lệnh có thể thực thi được bởi một hoặc nhiều bộ xử lý. Một hoặc nhiều bộ xử lý có thể được tạo cấu hình để thu nhận ảnh tham chiếu I được kết hợp với khối video trong tín hiệu video. Một hoặc nhiều bộ xử lý cũng có thể được tạo cấu hình để thu nhận các mẫu dự đoán $I(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu I . i và j biểu diễn tọa độ của một mẫu với khối video. Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để điều khiển các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng việc dịch chuyển sang phải cho các thông số PROF bên trong dựa trên giá trị dịch chuyển bit để đạt được sự chính xác được thiết lập trước. Các thông số PROF bên trong có thể gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu dự đoán $I(i, j)$. Một hoặc nhiều bộ xử lý cũng có thể được tạo cấu hình để thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên quy trình dẫn xuất PROF đang được áp dụng cho khối video dựa trên các mẫu dự đoán $I(i, j)$. Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để thu nhận các mẫu dự đoán của khối video dựa trên tổ hợp của các mẫu dự đoán và các giá trị tinh chỉnh dự đoán.

Theo khía cạnh thứ năm, sáng chế đề xuất thiết bị điện toán. Thiết bị điện toán có thể gồm một hoặc nhiều bộ xử lý, bộ nhớ đọc được bởi máy tính không chuyển tiếp lưu trữ các lệnh có thể thực thi được bởi một hoặc nhiều bộ xử lý. Một hoặc nhiều bộ xử lý

có thể được tạo cấu hình để thu nhận ảnh tham chiếu thứ nhất $I^{(0)}$ và ảnh tham chiếu thứ hai $I^{(1)}$ được kết hợp với khối video. Ảnh tham chiếu thứ nhất $I^{(0)}$ ở trước ảnh hiện tại và ảnh tham chiếu thứ hai $I^{(1)}$ ở sau ảnh hiện tại trong thứ tự hiển thị. Một hoặc nhiều bộ xử lý có thể được tạo cấu hình để thu nhận các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ nhất $I^{(0)}$. i và j có thể biểu diễn tọa độ của một mẫu với ảnh hiện tại. Một hoặc nhiều bộ xử lý cũng có thể được tạo cấu hình để thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên quy trình dẫn xuất PROF đang được áp dụng cho khối video dựa trên các mẫu dự đoán $I^{(0)}(i, j)$. Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để thu nhận các mẫu dự đoán thứ hai $I^{(1)}(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ hai $I^{(1)}$. Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để điều khiển các thông số BDOF bên trong của quy trình dẫn xuất BDOF bằng cách áp dụng việc dịch chuyển cho các thông số BDOF bên trong. Các thông số BDOF bên trong bao gồm các giá trị gradient ngang và các giá trị gradient dọc được dẫn xuất dựa trên các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$, các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, các sự chênh lệch mẫu giữa các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ và các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, và các thông số dẫn xuất BDOF trung gian. Các thông số dẫn xuất BDOF trung gian có thể gồm các thông số sGxdI, sGydI, sGx2, sGxGy, và sGy2. sGxdI và sGydI có thể gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị chênh lệch mẫu, và giữa các giá trị gradient dọc và các giá trị chênh lệch mẫu. sGx2 và sGy2 có thể gồm các giá trị tương liên tự động của các giá trị gradient ngang và các giá trị gradient dọc. sGxGy có thể gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị gradient dọc. Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để thu nhận các sự tinh chỉnh chuyển động cho các mẫu trong khối video dựa trên BDOF đang được áp dụng cho khối video dựa trên các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ và các mẫu dự đoán thứ hai $I^{(1)}(i, j)$.

Một hoặc nhiều bộ xử lý còn có thể được tạo cấu hình để thu nhận các mẫu dự đoán đôi của khối video dựa trên các sự tinh chỉnh chuyển động.

Theo khía cạnh thứ sáu, sáng chế đề xuất phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp đã lưu trữ các lệnh ở trong nó. Khi các lệnh được thực thi bởi một hoặc nhiều bộ xử lý của thiết bị, thì các lệnh có thể làm cho thiết bị nhận ba cờ điều khiển trong tập hợp thông số chuỗi (SPS). Cờ điều khiển thứ nhất chỉ ra liệu BDOF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ hai chỉ ra liệu PROF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ ba chỉ ra liệu DMVR có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Các lệnh cũng có thể làm cho thiết bị nhận cờ có mặt thứ nhất trong SPS khi cờ điều khiển thứ nhất là đúng, cờ có mặt thứ hai trong SPS khi cờ điều khiển thứ hai là đúng và cờ có mặt thứ ba trong SPS khi cờ điều khiển thứ ba là đúng. Các lệnh còn có thể làm cho thiết bị nhận cờ điều khiển ảnh thứ nhất trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ nhất trong SPS chỉ ra rằng BDOF không được cho phép cho các khối video ở trong ảnh. Các lệnh cũng có thể làm cho thiết bị nhận cờ điều khiển ảnh thứ hai trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ hai trong SPS chỉ ra rằng PROF không được cho phép cho các khối video ở trong ảnh. Các lệnh còn có thể làm cho thiết bị nhận cờ điều khiển ảnh thứ ba trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ ba trong SPS chỉ ra rằng DMVR không được cho phép cho các khối video ở trong ảnh.

Cần hiểu rằng cả phần mô tả tổng quát nêu trên lẫn phần mô tả chi tiết sau đây là chỉ được chỉ là các ví dụ và không nhằm để giới hạn sáng chế.

Mô tả vắn tắt các hình vẽ

Các hình vẽ kèm theo, vốn được tích hợp ở đây, tạo thành một phần của bản mô tả này, minh họa các khía cạnh làm ví dụ của sáng chế này và cùng với phần mô tả, có tác dụng giải thích nguyên lý của sáng chế.

Fig.1 là sơ đồ khối của bộ mã hóa, theo ví dụ của sáng chế.

Fig.2 là sơ đồ khối của bộ giải mã, theo ví dụ của sáng chế.

Fig.3A là sơ đồ minh họa các sự phân chia khối trong cấu trúc cây nhiều loại, theo ví dụ của sáng chế.

Fig.3B là sơ đồ minh họa các sự phân chia khối trong cấu trúc cây nhiều loại, theo ví dụ của sáng chế.

Fig.3C là sơ đồ minh họa các sự phân chia khối trong cấu trúc cây nhiều loại, theo ví dụ của sáng chế.

Fig.3D là sơ đồ minh họa các sự phân chia khối trong cấu trúc cây nhiều loại, theo ví dụ của sáng chế.

Fig.3E là sơ đồ minh họa các sự phân chia khối trong cấu trúc cây nhiều loại, theo ví dụ của sáng chế.

Fig.4 là sự minh họa dạng sơ đồ của mô hình luồng quang học hai chiều (BD OF), theo ví dụ của sáng chế.

Fig.5A là sự minh họa của mô hình afin, theo ví dụ của sáng chế.

Fig.5B là sự minh họa của mô hình afin, theo ví dụ của sáng chế.

Fig.6 là sự minh họa của mô hình afin, theo ví dụ của sáng chế.

Fig.7 là sự minh họa của sự tinh chỉnh dự đoán với luồng quang học (PROF), theo ví dụ của sáng chế.

Fig.8 là tiến trình công việc của BDOF, theo ví dụ của sáng chế.

Fig.9 là tiến trình công việc của PROF, theo ví dụ của sáng chế.

Fig.10 là phương pháp của BDOF, theo ví dụ của sáng chế.

Fig.11 là phương pháp của BDOF và PROF, theo ví dụ của sáng chế.

Fig.12 là phương pháp của BDOF, PROF, và DMVR, theo ví dụ của sáng chế.

Fig.13 là sự minh họa của tiến trình công việc của PROF cho việc dự đoán đôi, theo ví dụ của sáng chế.

Fig.14 là sự minh họa của các giai đoạn đường ống của quy trình BDOF và PROF, theo sáng chế.

Fig.15 là sự minh họa của phương pháp dẫn xuất gradient của BDOF, theo sáng chế.

Fig.16 là sự minh họa của phương pháp dẫn xuất gradient của PROF, theo sáng chế.

Fig.17A là sự minh họa của việc dẫn xuất các mẫu khuôn mẫu cho chế độ afin, theo ví dụ của sáng chế.

Fig.17B là sự minh họa của việc dẫn xuất các mẫu khuôn mẫu cho chế độ afin, theo ví dụ của sáng chế.

Fig.18A là sự minh họa của việc cho phép theo cách dành riêng PROF và LIC cho chế độ afin, theo ví dụ của sáng chế.

Fig.18B là sự minh họa của cho phép theo cách chung PROF và LIC cho chế độ afin, theo ví dụ của sáng chế.

Fig.19A là sơ đồ minh họa phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, theo ví dụ của sáng chế.

Fig.19B là sơ đồ minh họa phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, theo ví dụ của sáng chế.

Fig.19C là sơ đồ minh họa phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, theo ví dụ của sáng chế.

Fig.19D là sơ đồ minh họa phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, theo ví dụ của sáng chế.

Fig.20 là sơ đồ minh họa môi trường điện toán được ghép nối với giao diện người dùng, theo ví dụ của sáng chế.

Mô tả chi tiết sáng chế

Bây giờ, sáng chế sẽ được mô tả chi tiết đến các phương án khác nhau, ví dụ về các phương án này được minh họa trong các hình vẽ kèm theo. Phần mô tả sau đây đề cập đến các hình vẽ kèm theo mà trong đó các số chỉ dẫn giống nhau trong các hình vẽ khác nhau biểu diễn các phần tử giống hoặc tương tự nhau trừ khi được biểu diễn theo cách khác. Các cách triển khai được nêu trong phần mô tả sau đây của các phương án không biểu diễn tất cả các cách triển khai phù hợp với sáng chế. Thay vào đó, chúng chỉ là các ví dụ về các thiết bị và các phương pháp phù hợp với các khía cạnh liên quan đến sáng chế được nêu trong các yêu cầu bảo hộ kèm theo.

Hệ thuật ngữ được sử dụng trong sáng chế chỉ nhằm mục đích mô tả các phương án cụ thể và không được dự tính để giới hạn sáng chế. Như được sử dụng trong sáng chế

và các yêu cầu bảo hộ kèm theo, các dạng số ít cũng được dự tính để gồm các dạng số nhiều, trừ khi được chỉ ra một cách rõ ràng theo cách khác. Cũng cần hiểu rằng, thuật ngữ “và/hoặc” được sử dụng ở đây để chỉ và gồm bất kỳ hoặc tất cả các tổ hợp khả thi của một hoặc nhiều mục được liệt kê liên quan.

Cần hiểu rằng, mặc dù các thuật ngữ “thứ nhất”, “thứ hai”, “thứ ba”, v.v. có thể được sử dụng ở đây để mô tả các thông tin khác nhau, nhưng không tin không bị giới hạn ở các thuật ngữ này. Các thuật ngữ này chỉ được sử dụng để phân biệt một phân loại của thông tin so với thông tin khác. Ví dụ, không vượt ra khỏi phạm vi của sáng chế này, thông tin thứ nhất có thể được đặt là thông tin thứ hai, và tương tự, thông tin thứ hai cũng có thể được đặt là thông tin thứ nhất. Như được sử dụng ở đây, thuật ngữ “nếu” có thể được hiểu nghĩa là “khi” hoặc “lúc” hoặc “để đáp lại một phán quyết”, phụ thuộc vào ngữ cảnh.

Phiên bản thứ nhất của tiêu chuẩn HEVC được hoàn thành trong tháng 10 năm 2014, đề xuất việc tiết kiệm xấp xỉ 50% tỉ lệ bit hoặc chất lượng nhận thức tương tự khi so với tiêu chuẩn mã hóa video thế hệ trước đó H.264/MPEG AVC. Mặc dù tiêu chuẩn HEVC cung cấp các sự cải thiện tạo mã đáng kể hơn so với các tiêu chuẩn tiền nhiệm của nó, nhưng có bằng chứng là có thể để đạt được hiệu quả tạo mã vượt trội hơn HEVC, với các công cụ tạo mã bổ sung. Dựa trên điều đó, cả VCEG và MPEG đều đã bắt đầu công việc khám phá các công nghệ tạo mã mới cho việc chuẩn hóa tạo mã video trong tương lai. Một nhóm khám phá video hợp tác (Joint Video Exploration Team, JVET) đã được tạo thành vào tháng 10 năm 2015 bởi ITU-T VCEG và ISO/IEC MPEG để bắt đầu nghiên cứu đáng kể về các công nghệ cải tiến mà có thể cho phép nâng cao đáng kể hiệu quả tạo mã. Một phần mềm tham chiếu được gọi là mô hình khám phá hợp tác (Joint

Exploration Model, JEM) đã được duy trì bởi JVET bằng cách tích hợp nhiều công cụ tạo mã bổ sung trên đỉnh của mô hình kiểm tra HEVC (HEVC test Model, HM).

Vào tháng 10 năm 2017, lời kêu gọi hợp tác cho các đề xuất (Call for Proposal, CfP) về việc nén video với khả năng vượt qua HEVC đã được phát hành bởi ITU-T và ISO/IEC **Error! Reference source not found.** Vào tháng 4 năm 2018, 23 phản hồi CfP được nhận và được đánh giá tại phiên họp JVET lần thứ 10, đã chứng minh rằng hiệu quả nén đã vượt quá HEVC khoảng quanh 40%. Dựa trên các kết quả đánh giá này, JVET đã triển khai dự án mới để phát triển tiêu chuẩn tạo mã video thế hệ mới có tên là tạo mã video vạn năng (Versatile Video Coding, VVC). Trong cùng tháng đó, một cơ sở mã phần mềm tham chiếu, được gọi là mô hình kiểm tra VVC (VVC Test Model, VTM), được thiết lập để chứng minh cách triển khai tham chiếu của tiêu chuẩn VVC.

Giống như HEVC, VVC được xây dựng trên khung tạo mã video lai dựa trên khối.

Fig.1 thể hiện sơ đồ chung của bộ mã hóa video dựa trên khối cho VVC. Cụ thể là, Fig.1 thể hiện bộ mã hóa điển hình 100. Bộ mã hóa 100 có đầu vào video 110, việc bù chuyển động 112, việc ước lượng chuyển động 114, việc quyết định chế độ trong ảnh/liên ảnh 116, bộ dự đoán khối 140, bộ cộng 128, việc biến đổi 130, việc lượng tử hóa 132, thông tin liên quan đến việc dự đoán 142, việc dự đoán trong ảnh 118, bộ đệm ảnh 120, việc lượng tử hóa ngược 134, việc biến đổi ngược 136, bộ cộng 126, bộ nhớ 124, bộ lọc trong vòng 122, tạo mã entropi 138, và dòng bit 144.

Trong bộ mã hóa 100, khung video được phân chia thành nhiều khối video để xử lý. Đối với mỗi khối video đã cho, việc dự đoán được tạo thành dựa trên hoặc là cách tiếp cận dự đoán liên ảnh hoặc là cách tiếp cận dự đoán trong ảnh.

Phần dư dự đoán, việc biểu diễn sự chênh lệch giữa khối video hiện tại, một phần của đầu vào video 110, và bộ dự đoán của nó, một phần của bộ dự đoán khối 140, được

gửi đến việc biến đổi 130 từ bộ cộng 128. Sau đó, các hệ số biến đổi được gửi từ việc biến đổi 130 tới việc lượng tử hóa 132 để giảm entropi. Sau đó, các hệ số được lượng tử hóa được nạp vào việc tạo mã entropi 138 để tạo ra dòng bit video được nén. Như được thể hiện trên Fig.1, thông tin liên quan đến việc dự đoán 142 từ việc quyết định chế độ trong ảnh/liên ảnh 116, như là thông tin phân chia khối video, các vectơ chuyển động, chỉ số ảnh tham chiếu, và chế độ dự đoán trong ảnh, cũng được nạp thông qua hệ mạch tạo mã entropi 138 và được lưu vào trong dòng bit video được nén 144. Dòng bit được nén 144 gồm dòng bit video.

Trong bộ mã hóa 100, các hệ mạch liên quan đến bộ giải mã cũng là cần thiết để tái tạo các điểm ảnh cho mục đích dự đoán. Trước tiên, phần dư dự đoán được tái tạo qua việc lượng tử hóa ngược 134 và việc biến đổi ngược 136. Phần dư dự đoán được tái tạo này được tổ hợp với bộ dự đoán khối 140 để tạo ra các điểm ảnh được tái tạo mà chưa được lọc cho khối video hiện tại.

Việc dự đoán theo không gian (hoặc “việc dự đoán trong ảnh”) sử dụng các điểm ảnh từ các mẫu của các khối lân cận hiện đã được tạo mã (vốn được gọi là các mẫu tham chiếu) trong cùng một khung video như video hiện tại để dự đoán khối video hiện tại.

Việc dự đoán theo thời gian (còn được gọi là “việc dự đoán liên ảnh”) sử dụng các điểm ảnh được tái tạo từ các ảnh video hiện đã được tạo mã để dự đoán khối video hiện tại. Việc dự đoán theo thời gian làm giảm tính dư thừa theo thời gian có hữu trong tín hiệu video. Tín hiệu dự đoán theo thời gian cho đơn vị tạo mã (Coding Unit, CU) hoặc khối tạo mã đã cho thường được phát tín hiệu bởi một hoặc nhiều MV vốn chỉ ra lượng và hướng của chuyển động giữa CU hiện tại và tham chiếu theo thời gian của nó. Hơn nữa, nếu nhiều ảnh tham chiếu được hỗ trợ, thì một chỉ số ảnh tham chiếu sẽ được gửi

theo cách bổ sung, vốn được sử dụng để nhận dạng nơi mà ảnh tham chiếu ở trong ảnh tham chiếu lưu trữ tín hiệu dự đoán theo thời gian, đi ra từ đó.

Việc ước lượng chuyển động 114 lấy đầu vào video 110 và tín hiệu từ bộ đệm ảnh 120 và được xuất ra, đến việc bù chuyển động 112, tín hiệu ước lượng chuyển động. Việc bù chuyển động 112 lấy đầu vào video 110, tín hiệu từ bộ đệm ảnh 120, và tín hiệu ước lượng chuyển động từ việc ước lượng chuyển động 114 và được xuất ra đến việc quyết định chế độ trong ảnh/liên ảnh 116, tín hiệu bù chuyển động.

Sau khi việc dự đoán theo không gian và/hoặc theo thời gian được thực hiện, thì việc quyết định chế độ 116 trong bộ mã hóa 100 chọn chế độ dự đoán tốt nhất, ví dụ, dựa trên phương pháp tối ưu hóa nhiễu loạn tỉ lệ. Sau đó, khối dự đoán 140 được trừ khỏi khối video hiện tại và phần dư dự đoán được tạo thành được giải tương liên nhờ sử dụng việc biến đổi 130 và việc lượng tử hóa 132. Các hệ số dư được lượng tử hóa kết quả sẽ được lượng tử hóa ngược bởi việc lượng tử hóa ngược 134 và biến đổi ngược bởi việc biến đổi ngược 136 để tạo thành phần dư được tái tạo, sau đó được cộng lại vào khối dự đoán để tạo thành tín hiệu được tái tạo của CU. Việc lọc trong vòng 122 thêm nữa, như là bộ lọc khử khối, dịch vị thích ứng mẫu (Sample Adaptive Offset, SAO), và/hoặc bộ lọc vòng thích ứng (Adaptive in-Loop Filter, ALF) có thể được áp dụng trên CU được tái tạo trước khi nó được nhập trong bộ phận lưu trữ ảnh tham chiếu của bộ đệm ảnh 120 và được sử dụng để các khối tạo mã video tương lai. Để tạo thành dòng bit video đầu ra 144, thì chế độ mã hóa (liên ảnh hoặc trong ảnh), thông tin chế độ dự đoán, thông tin chuyển động, và các hệ số dư được lượng tử hóa đều được gửi đến đơn vị tạo mã entropi 138 để được nén thêm nữa và đóng gói để tạo thành dòng bit.

Fig.1 đưa ra sơ đồ khối của hệ thống mã hóa video lai dựa trên khối chung. Tín hiệu video đầu vào được xử lý bởi khối (được gọi là các CU). Trong VTM-1.0, CU có

thể lên tới 128x128 điểm ảnh. Tuy nhiên, khác với HEVC vốn phân chia các khối chỉ dựa trên các cây tứ phân, thì trong VVC, một đơn vị cây tạo mã (Coding Tree Unit, CTU) được chia tách thành các CU để thích ứng để biến đổi các đặc tính cục bộ dựa trên cây tứ phân/nhị phân/tam phân. Theo cách bổ sung, khái niệm của loại đơn vị nhiều sự phân chia trong HEVC được xóa bỏ, tức là, việc tách ra của CU, đơn vị dự đoán (Prediction Unit, PU) và đơn vị biến đổi (Transform Unit, TU) không tồn tại trong VVC nữa; thay vào đó, mỗi CU luôn được sử dụng như là đơn vị cơ bản cho cả việc dự đoán và biến đổi mà không có thêm các sự phân chia thêm nữa. Trong cấu trúc cây nhiều loại, một CTU được phân chia trước tiên bởi cấu trúc cây tứ phân. Sau đó, mỗi nút lá cây tứ phân có thể được phân chia thêm nữa bởi cấu trúc cây nhị phân hoặc tam phân. Như được thể hiện trên Fig.3A, Fig.3B, Fig.3C, Fig.3D, và Fig.3E, thì có năm loại chia tách: phân chia tứ phân, phân chia nhị phân dọc, phân chia nhị phân ngang, phân chia tam phân ngang, và phân chia tam phân dọc.

Fig.3A thể hiện sơ đồ minh họa sự phân chia tứ phân khối trong cấu trúc cây nhiều loại, theo sáng chế.

Fig.3B thể hiện sơ đồ minh họa sự phân chia nhị phân dọc khối trong cấu trúc cây nhiều loại, theo sáng chế.

Fig.3C thể hiện sơ đồ minh họa sự phân chia nhị phân ngang khối trong cấu trúc cây nhiều loại, theo sáng chế.

Fig.3D thể hiện sơ đồ minh họa sự phân chia tam phân dọc khối trong cấu trúc cây nhiều loại, theo sáng chế.

Fig.3E thể hiện sơ đồ minh họa sự phân chia tam phân ngang khối trong cấu trúc cây nhiều loại, theo sáng chế.

Trên Fig.1, việc dự đoán theo không gian và/hoặc việc dự đoán theo thời gian có thể được thực hiện. Việc dự đoán theo không gian (hoặc “việc dự đoán trong ảnh”) sử dụng các điểm ảnh từ các mẫu của các khối lân cận hiện đã được mã hóa (vốn được gọi là các mẫu tham chiếu) trong cùng một hình ảnh/phiên video để dự đoán khối video hiện tại. Việc dự đoán theo không gian làm giảm tính dư thừa về mặt không gian cố hữu trong tín hiệu video. Việc dự đoán theo thời gian (còn được gọi là “việc dự đoán liên ảnh” hoặc “việc dự đoán được bù chuyển động”) sử dụng các điểm ảnh được tái tạo từ các ảnh video đã được tạo mã để dự đoán khối video hiện tại. Việc dự đoán theo thời gian làm giảm tính dư thừa theo thời gian cố hữu trong tín hiệu video. Tín hiệu dự đoán theo thời gian cho CU đã cho thường được phát tín hiệu bởi một hoặc nhiều vectơ chuyển động (Motion Vector, MV) vốn chỉ ra số lượng và hướng chuyển động giữa CU hiện tại và tham chiếu theo thời gian của nó. Ngoài ra, nếu nhiều hình ảnh tham chiếu được hỗ trợ, thì một chỉ số hình ảnh tham chiếu sẽ được gửi theo cách bổ sung, vốn được sử dụng để nhận dạng nơi mà hình ảnh tham chiếu trong hình ảnh tham chiếu lưu trữ tín hiệu dự đoán theo thời gian, đi ra từ đó. Sau việc dự đoán theo không gian và/hoặc theo thời gian, thì khối quyết định chế độ trong bộ mã hóa chọn chế độ dự đoán tốt nhất, ví dụ, dựa trên phương pháp tối ưu hóa nhiễu loạn tỉ lệ. Sau đó, khối dự đoán được trừ khỏi khối video hiện tại và phần dư dự đoán được giải tương liên nhờ sử dụng sự biến đổi và được lượng tử hóa. Các hệ số dư được lượng tử hóa sẽ được lượng tử hóa ngược và biến đổi ngược để tạo thành phần dư được tái tạo, mà sau đó được cộng vào khối dự đoán để tạo thành tín hiệu được tái tạo của CU. Hơn nữa, việc lọc trong vòng, như bộ lọc khử khối, dịch vị tương thích mẫu (Sample Adaptive Offset, SAO), và/hoặc bộ lọc trong vòng thích ứng (Adaptive in-Loop Filter, ALF) có thể được áp dụng trên CU được tái tạo trước khi nó được nhập trong bộ lưu trữ ảnh tham chiếu và được sử dụng để các khối tạo mã video tương lai. Để tạo thành dòng bit video đầu ra, thì chế độ mã hóa (liên ảnh

hoặc trong ảnh), thì thông tin chế độ dự đoán, thông tin chuyển động, và các hệ số dư được lượng tử hóa đều được gửi đến đơn vị tạo mã entropi để còn được nén và đóng gói để tạo thành dòng bit.

Fig.2 thể hiện sơ đồ khối chung của bộ giải mã video cho VVC. Cụ thể là, Fig.2 thể hiện sơ đồ khối bộ giải mã điển hình 200. Bộ giải mã 200 có dòng bit 210, việc giải mã entropi 212, việc lượng tử hóa ngược 214, việc biến đổi ngược 216, bộ cộng 218, việc lựa chọn chế độ trong ảnh/liên ảnh 220, việc dự đoán trong ảnh 222, bộ nhớ 230, bộ lọc trong vòng 228, việc bù chuyển động 224, bộ đệm ảnh 226, thông tin liên quan đến việc dự đoán 234, và đầu ra video 232.

Bộ giải mã 200 là tương tự như phần liên quan đến việc tái tạo nằm trong bộ mã hóa 100 trên Fig.1. Trong bộ giải mã 200, dòng bit video đi tới 210 được giải mã trước tiên thông qua việc giải mã entropi 212 để dẫn xuất các mức hệ số được lượng tử hóa và thông tin liên quan đến việc dự đoán. Sau đó, các mức hệ số được lượng tử hóa được xử lý thông qua việc lượng tử hóa ngược 214 và việc biến đổi ngược 216 để thu nhận phần dư dự đoán được tái tạo. Cơ chế bộ dự đoán khối, được triển khai trong bộ lựa chọn chế độ trong ảnh/liên ảnh 220, được tạo cấu hình để thực hiện hoặc là việc dự đoán trong ảnh 222, hoặc là việc bù chuyển động 224, dựa trên thông tin dự đoán được giải mã. Tập hợp của các điểm ảnh được tái tạo mà chưa được lọc sẽ được thu nhận bằng cách tính tổng phần dư dự đoán được tái tạo từ việc biến đổi ngược 216 và đầu ra dự đoán được tạo ra bởi cơ chế bộ dự đoán khối, nhờ sử dụng bộ tính tổng 218.

Khối được tái tạo có thể còn đi qua bộ lọc trong vòng 228 trước khi nó được lưu trữ trong bộ đệm ảnh 226, vốn có chức năng như là bộ phận lưu trữ ảnh tham chiếu. Sau đó, video được tái tạo trong bộ đệm ảnh 226 có thể được gửi đi để điều vận thiết bị hiển thị, cũng như được sử dụng để dự đoán các khối video tương lai. Trong các tình huống

mà tại đó bộ lọc trong vòng 228 được bật lên, thì hoạt động lọc được thực hiện trên các điểm ảnh được tái tạo này để dẫn xuất đầu ra video được tái tạo cuối cùng 232.

Fig.2 đưa ra sơ đồ khối chung của bộ giải mã video dựa trên khối. Dòng bit video được giải mã entropi trước tiên tại đơn vị giải mã entropi. Chế độ tạo mã và thông tin dự đoán được gửi đến hoặc là đơn vị dự đoán theo không gian (nếu được tạo mã trong ảnh) hoặc đến đơn vị dự đoán theo thời gian (nếu được tạo mã liên ảnh) để tạo thành khối dự đoán. Các hệ số biến đổi dư được gửi đến đơn vị lượng tử hóa ngược và đơn vị biến đổi ngược để tái tạo khối dư. Sau đó, khối dự đoán và khối dư được cộng lại với nhau. Khối được tái tạo có thể còn trải qua việc lọc trong vòng trước khi nó được lưu trữ trong bộ lưu trữ ảnh tham chiếu. Sau đó, video được tái tạo trong bộ lưu trữ ảnh tham chiếu được gửi ra để điều vận thiết bị hiển thị, cũng như được sử dụng để dự đoán các khối video tương lai.

Nói chung, các kỹ thuật dự đoán liên ảnh cơ bản mà được áp dụng trong VVC được giữ giống như các kỹ thuật dự đoán liên ảnh của HEVC ngoại trừ một vài mô đun được mở rộng và/hoặc được tăng cường thêm nữa. Một cách cụ thể, đối với tất cả các tiêu chuẩn video trước đây, một khối tạo mã có thể chỉ được kết hợp với một MV đơn lẻ khi khối tạo mã được dự đoán đơn hoặc hai MV khi khối tạo mã được dự đoán đôi. Bởi vì sự giới hạn như vậy của việc bù chuyển động dựa trên khối thông thường, nên chuyển động nhỏ vẫn có thể ở trong các mẫu dự đoán sau khi bù chuyển động, do đó ảnh hưởng một cách tiêu cực lên hiệu quả tổng thể của việc bù chuyển động. Để cải thiện cả độ chi tiết và độ chính xác của các MV, thì hai phương pháp tinh chỉnh đảo mẫu dựa trên luồng quang học, tức là luồng quang học hai chiều (BDOF) và sự tinh chỉnh dự đoán với luồng quang học (PROF) cho chế độ afin, hiện đang được nghiên cứu cho tiêu

chuẩn VVC. Trong phần sau đây, các khía cạnh kỹ thuật chính của hai công cụ tạo mã liên ảnh được xem xét một cách ngắn gọn.

Luồng quang học hai chiều

Trong VVC, BDOF được áp dụng để tinh chỉnh các mẫu dự đoán của các khối tạo mã được dự đoán đôi. Cụ thể là, như được thể hiện trên Fig.4, BDOF là sự tinh chỉnh chuyển động đảo mẫu được thực hiện trên đỉnh của các dự đoán được bù chuyển động dựa trên khối khi việc dự đoán đôi được sử dụng.

Fig.4 thể hiện sự minh họa của mô hình BDOF, theo sáng chế.

Sự tinh chỉnh chuyển động của mỗi khối con 4×4 được tính toán bằng cách tối thiểu hóa sự chênh lệch giữa các mẫu dự đoán L0 và L1 sau khi BDOF được áp dụng nằm trong một cửa sổ 6×6 Ω quanh khối con. Cụ thể là, giá trị của (v_x, v_y) được dẫn xuất dưới dạng

$$\begin{aligned} v_x &= S_1 > 0? \text{clip3} \left(-th_{BDOF}, th_{BDOF}, -((S_3 \cdot 2^3) \gg \lfloor \log_2 S_1 \rfloor) \right) : 0 \\ v_y &= S_5 > 0? \text{clip3} \left(-th_{BDOF}, th_{BDOF}, - \left((S_6 \cdot 2^3 \right. \right. \\ &\quad \left. \left. - \left((v_x S_{2,m}) \ll n_{S_2} + v_x S_{2,s} \right) / 2 \right) \gg \lfloor \log_2 S_5 \rfloor \right) : 0 \end{aligned} \quad (1)$$

mà tại đó $\lfloor \cdot \rfloor$ là hàm sàn; $\text{clip3}(\min, \max, x)$ là hàm mà cắt xén giá trị đã cho x bên trong phạm vi $[\min, \max]$; ký hiệu $\gg$ biểu diễn phép toán dịch dịch chuyển sang phải đảo bit; ký hiệu $\ll$ biểu diễn phép toán dịch chuyển sang trái đảo bit; th_{BDOF} là ngưỡng tinh chỉnh chuyển động để ngăn chặn các sai số truyền do chuyển động cục bộ không đều, vốn bằng $1 \ll \max(5, \text{bit depth} - 7)$, mà tại đó bit depth là độ sâu bit bên trong.

Trong (1), $S_{2,m} = S_2 \gg n_{S_2}$, $S_{2,s} = S_2 \& (2^{n_{S_2}} - 1)$.

Các giá trị S_1, S_2, S_3, S_5 và S_6 được tính toán dưới dạng

$$\begin{aligned}
S_1 &= \sum_{(i,j) \in \Omega} \psi_x(i,j) \cdot \psi_x(i,j) \\
S_2 &= \sum_{(i,j) \in \Omega} \psi_x(i,j) \cdot \psi_y(i,j) \\
S_3 &= \sum_{(i,j) \in \Omega} \theta(i,j) \cdot \psi_x(i,j) \\
S_5 &= \sum_{(i,j) \in \Omega} \psi_y(i,j) \cdot \psi_y(i,j) \\
S_6 &= \sum_{(i,j) \in \Omega} \theta(i,j) \cdot \psi_y(i,j)
\end{aligned} \tag{2}$$

mà tại đó

$$\begin{aligned}
\psi_x(i,j) &= \left(\frac{\partial I^{(1)}}{\partial x}(i,j) + \frac{\partial I^{(0)}}{\partial x}(i,j) \right) \gg \max(1, \text{bit depth} - 11) \\
\psi_y(i,j) &= \left(\frac{\partial I^{(1)}}{\partial y}(i,j) + \frac{\partial I^{(0)}}{\partial y}(i,j) \right) \gg \max(1, \text{bitdepth} - 11) \\
\theta(i,j) &= (I^{(1)}(i,j) \gg \max(4, \text{bitdepth} - 8)) \\
&\quad - (I^{(0)}(i,j) \gg \max(4, \text{bitdepth} - 8))
\end{aligned} \tag{3}$$

mà tại đó $I^{(k)}(i,j)$ là giá trị mẫu tại tọa độ (i,j) của tín hiệu dự đoán trong danh

sách $k, k = 0,1$, vốn được tạo ra tại sự chính xác cao trung gian (tức là, 16 bit); $\frac{\partial I^{(k)}}{\partial x}(i,j)$

và $\frac{\partial I^{(k)}}{\partial y}(i,j)$ là các gradient ngang và dọc của mẫu mà được thu nhận bằng cách tính toán

một cách trực tiếp sự chênh lệch giữa hai mẫu lân cận của nó, tức là,

$$\begin{aligned}
\frac{\partial I^{(k)}}{\partial x}(i,j) &= (I^{(k)}(i+1,j) - I^{(k)}(i-1,j)) \gg \max(6, \text{bit-depth} - 6) \\
\frac{\partial I^{(k)}}{\partial y}(i,j) &= (I^{(k)}(i,j+1) - I^{(k)}(i,j-1)) \gg \max(6, \text{bit-depth} - 6)
\end{aligned} \tag{4}$$

Dựa trên sự tinh chỉnh chuyển động được dẫn xuất trong (1(1)), thì các mẫu dự đoán hai chiều cuối cùng của CU được tính toán bằng cách ngoại suy các mẫu dự đoán L0/L1 dọc theo quỹ đạo chuyển động dựa trên mô hình luồng quang, như được chỉ ra bởi

$$\begin{aligned}
 pred_{BDOF}(x,y) &= (I^{(0)}(x,y) + I^{(1)}(x,y) + b + o_{offset}) \gg shift \\
 &+ b \\
 &= rnd\left(\left(v_x\left(\frac{\partial I^{(1)}(x,y)}{\partial x} - \frac{\partial I^{(0)}(x,y)}{\partial x}\right)\right)/2\right) + rnd \\
 &\quad \left(\left(v_y\left(\frac{\partial I^{(1)}(x,y)}{\partial y} - \frac{\partial I^{(0)}(x,y)}{\partial y}\right)\right)/2\right)
 \end{aligned} \tag{5}$$

mà tại đó $shift$ và o_{offset} là giá trị dịch chuyển sang phải và giá trị dịch vị mà được áp dụng để tổ hợp các tín hiệu dự đoán L0 và L1 cho việc dự đoán đôi, vốn bằng với $15 - bitdepth$ và $1 \ll (14 - bitdepth) + 2 \cdot (1 \ll 13)$, một cách tương ứng. Dựa trên phương pháp điều khiển độ sâu bit ở trên, thì bảo đảm được rằng độ sâu bit tối đa của các thông số trung gian của toàn bộ quy trình BDOF không vượt quá 32-bit và đầu vào lớn nhất cho phép nhân là ở trong 15-bit, tức là, một số nhân 15-bit là đủ cho các cách triển khai BDOF.

Chế độ affin

Trong HEVC, thì chỉ mô hình chuyển động tịnh tiến được áp dụng cho việc dự đoán được bù chuyển động. Trong khi trong thế giới thực, thì có nhiều kiểu chuyển động, ví dụ, các chuyển động phóng to/thu nhỏ, quay, phối cảnh và các chuyển động không đều khác. Trong VVC, việc dự đoán được bù chuyển động affin được áp dụng bằng cách phát tín hiệu một cờ cho mỗi khối tạo mã liên ảnh để chỉ ra liệu chuyển động tịnh tiến hoặc mô hình chuyển động affin được áp dụng cho việc dự đoán liên ảnh. Trong thiết kế

VVC hiện tại, thì hai chế độ afin, gồm chế độ afin 4-thông số và chế độ afin 6-thông số, được hỗ trợ cho một khối tạo mã afin.

Mô hình afin 4-thông số có các thông số sau đây: hai thông số cho sự di chuyển tịnh tiến theo các hướng ngang và dọc một cách tương ứng, một thông số cho chuyển động thu phóng và một thông số cho chuyển động quay theo cả hai hướng. Thông số thu phóng ngang là bằng thông số thu phóng dọc. Thông số quay ngang là bằng thông số quay dọc. Để đạt được sự tương thích tốt nhất của các vectơ chuyển động và thông số afin, trong VVC, thì các thông số afin đó được tịnh tiến thành hai MV (cũng được gọi là vectơ chuyển động điểm điều khiển (Control Point Motion Vector, CPMV)) được đặt tại góc trên cùng bên trái và góc trên cùng bên phải của khối hiện tại. Như được thể hiện trên Fig.5A và Fig.5B, trường chuyển động afin của khối được mô tả bởi hai MV điểm điều khiển (V_0, V_1).

Fig.5A thể hiện sự minh họa của mô hình afin 4-thông số, theo sáng chế.

Fig.5B thể hiện sự minh họa của mô hình afin 4-thông số, theo sáng chế.

Dựa trên chuyển động điểm điều khiển, thì trường chuyển động (v_x, v_y) của một khối được tạo mã afin được mô tả dưới dạng

$$\begin{aligned} v_x &= \frac{(v_{1x} - v_{0x})}{w} x - \frac{(v_{1y} - v_{0y})}{w} y + v_{0x} \\ v_y &= \frac{(v_{1y} - v_{0y})}{w} x + \frac{(v_{1x} - v_{0x})}{w} y + v_{0y} \end{aligned} \quad (6)$$

Chế độ afin 6-thông số có các thông số sau đây: hai thông số cho sự di chuyển tịnh tiến theo các hướng ngang và dọc một cách tương ứng, một thông số cho chuyển động thu phóng và một thông số cho chuyển động quay theo hướng ngang, một thông số cho chuyển động thu phóng và một thông số cho chuyển động quay theo hướng dọc. Mô hình chuyển động afin 6-thông số được tạo mã với ba MV tại ba CPMV.

Fig.6 thể hiện sự minh họa của mô hình afin 6-thông số, theo sáng chế.

Như được thể hiện trên Fig.6, ba điểm điều khiển của một khối afin 6-thông số được đặt tại góc trên cùng bên trái, trên cùng bên phải, và dưới cùng bên trái của khối. Chuyển động tại điểm điều khiển trên cùng bên trái liên quan đến chuyển động tịnh tiến, và chuyển động tại điểm điều khiển trên cùng bên phải liên quan đến chuyển động quay và thu phóng theo hướng ngang, và chuyển động tại dưới cùng bên trái điểm điều khiển liên quan đến chuyển động quay và chuyển động thu phóng theo chiều dọc. So với mô hình chuyển động afin 4-thông số, thì chuyển động quay và thu phóng theo hướng ngang của 6-thông số có thể không giống với chuyển động quay và thu phóng theo hướng dọc. Giả sử (V_0, V_1, V_2) là các MV của các góc trên cùng bên trái, trên cùng bên phải và dưới cùng bên trái của khối hiện tại trên Fig.6, thì vectơ chuyển động của mỗi khối con (v_x, v_y) được dẫn xuất nhờ sử dụng ba MV tại các điểm điều khiển là:

$$\begin{aligned} v_x &= v_{0x} + (v_{1x} - v_{0x}) * \frac{x}{w} + (v_{2x} - v_{0x}) * \frac{y}{h} \\ v_y &= v_{0y} + (v_{1y} - v_{0y}) * \frac{x}{w} + (v_{2y} - v_{0y}) * \frac{y}{h} \end{aligned} \quad (7)$$

Sự tinh chỉnh dự đoán với luồng quang học cho chế độ afin

Để cải thiện độ chính xác bù chuyển động afin, thì PROF hiện đang được nghiên cứu trong VVC hiện tại vốn tinh chỉnh việc bù chuyển động afin dựa trên khối con dựa trên mô hình luồng quang học. Cụ thể là, sau khi thực hiện việc bù chuyển động afin dựa trên khối con, thì mẫu dự đoán độ chói của một khối afin được sửa đổi bởi một giá trị tinh chỉnh mẫu được dẫn xuất dựa trên phương trình luồng quang học. Chi tiết là, các phép toán của PROF có thể được tóm tắt theo bốn bước sau đây:

Bước một: Việc bù chuyển động affin dựa trên khối con được thực hiện để tạo ra khối dự đoán con $I(i, j)$ nhờ sử dụng các MV khối con như được dẫn xuất trong (6) đối với mô hình affin 4-thông số và (7) đối với mô hình affin 6-thông số.

Bước hai: Các gradient theo không gian $g_x(i, j)$ và $g_y(i, j)$ của mỗi mẫu dự đoán trong số các mẫu dự đoán được tính toán là

$$\begin{aligned} g_x(i, j) &= (I(i+1, j) - I(i-1, j)) \gg (\max(2, 14 - \text{bitdepth}) - 4) \\ g_y(i, j) &= (I(i, j+1) - I(i, j-1)) \gg (\max(2, 14 - \text{bitdepth}) - 4) \end{aligned} \quad (8)$$

Để tính toán các gradient, thì một hàng/cột bổ sung của các mẫu dự đoán cần được tạo ra trên mỗi phía của một khối con. Để giảm băng thông và độ phức tạp bộ nhớ, thì các mẫu trên các biên được mở rộng được sao chép từ vị trí điểm ảnh số nguyên gần nhất trong ảnh tham chiếu để tránh các quy trình nội suy bổ sung.

Bước ba: Giá trị tinh chỉnh dự đoán độ chói được tính toán bởi

$$\Delta I(i, j) = g_x(i, j) * \Delta v_x(i, j) + g_y(i, j) * \Delta v_y(i, j) \quad (9)$$

mà tại đó $\Delta v(i, j)$ là sự chênh lệch giữa điểm ảnh MV được tính cho sự định vị mẫu (i, j) , được miêu tả bởi $v(i, j)$, và MV khối con của khối con mà điểm ảnh (i, j) đặt tại đó. Theo cách bổ sung, trong thiết kế PROF hiện tại, sau khi cộng sự tinh chỉnh dự đoán vào mẫu dự đoán gốc, thì một phép toán cắt xén được thực hiện để cắt xén giá trị của mẫu dự đoán được tinh chỉnh là ở trong 15-bit, tức là,

$$\begin{aligned} I^r(i, j) &= I(i, j) + \Delta I(i, j) \\ I^r(i, j) &= \text{clip3}(-2^{14}, 2^{14} - 1, I^r(i, j)) \end{aligned}$$

mà tại đó $I(i, j)$ và $I^r(i, j)$ là mẫu dự đoán gốc và được tinh chỉnh tại sự định vị (i, j) , một cách tương ứng.

Fig.7 minh họa quy trình PROF cho chế độ afin, theo sáng chế. Fig.7 gồm khối 710, khối 720, và khối 730. Khối 730 là khối được quay của khối 720.

Bởi vì các thông số mô hình afin và sự định vị điểm ảnh tương đối với trung tâm khối con sẽ không được thay đổi từ khối con sang khối con, nên $\Delta v(i,j)$ có thể được tính toán cho khối con thứ nhất, và được sử dụng lại cho các khối con khác trong cùng CU. Đặt Δx và Δy là dịch vị ngang và dọc từ sự định vị mẫu (i,j) đến trung tâm của khối con mà mẫu thuộc về, thì $\Delta v(i,j)$ có thể được dẫn xuất dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= c * \Delta x + d * \Delta y \\ \Delta v_y(i,j) &= e * \Delta x + f * \Delta y\end{aligned}\tag{10}$$

Dựa trên các phương trình dẫn xuất MV khối con afin (6) và (7), thì sự chênh lệch MV có thể được dẫn xuất. Cụ thể là, đối với mô hình afin 4-thông số,

$$\begin{cases} c = f = \frac{v_{1x} - v_{0x}}{w} \\ e = -d = \frac{v_{1y} - v_{0y}}{w} \end{cases}$$

Đối với mô hình afin 6-thông số,

$$\begin{cases} c = \frac{v_{1x} - v_{0x}}{w} \\ d = \frac{v_{2x} - v_{0x}}{h} \\ e = \frac{v_{1y} - v_{0y}}{w} \\ f = \frac{v_{2y} - v_{0y}}{h} \end{cases}$$

mà tại đó (v_{0x}, v_{0y}) , (v_{1x}, v_{1y}) , (v_{2x}, v_{2y}) là các MV điểm điều khiển trên cùng bên trái, trên cùng bên phải và dưới cùng bên trái của khối tạo mã hiện tại, w và h là chiều rộng và chiều cao của khối. Trong thiết kế PROF sẵn có, thì sự chênh lệch MV Δv_x và Δv_y luôn được dẫn xuất tại độ chính xác là 1/32-pel.

Sự bù chiếu sáng cục bộ

Sự bù chiếu sáng cục bộ (Local Illumination Compensation, LIC) là công cụ tạo mã mà được sử dụng để giải quyết vấn đề về sự thay đổi chiếu sáng cục bộ tồn tại trong- giữa các ảnh lân cận theo thời gian. Một cặp các thông số trọng số và dịch vị được áp dụng cho các mẫu tham chiếu để thu nhận các mẫu dự đoán của một khối hiện tại. Mô hình toán học chung được đưa ra dưới dạng

$$P[\mathbf{x}] = \alpha * P_r[\mathbf{x} + \mathbf{v}] + \beta \quad (11)$$

mà tại đó $P_r[\mathbf{x} + \mathbf{v}]$ là khối tham chiếu được chỉ ra bởi véc-tơ chuyển động $\mathbf{v}$, $[\alpha, \beta]$ là một cặp các thông số trọng số và dịch vị tương ứng cho khối tham chiếu và $P[\mathbf{x}]$ là khối dự đoán cuối cùng. Một cặp các thông số trọng số và dịch vị được ước lượng nhờ sử dụng thuật toán sai số bình phương trung bình tuyến tính nhỏ nhất (Least Linear Mean Square Error, LLMSE) dựa trên khuôn mẫu (tức là, các mẫu được tái tạo lân cận) của khối hiện tại và khối tham chiếu của khuôn mẫu (vốn được dẫn xuất nhờ sử dụng véc-tơ chuyển động của khối hiện tại). Bằng cách tối thiểu hóa sự chênh lệch bình phương trung bình giữa các mẫu khuôn mẫu và các mẫu tham chiếu của khuôn mẫu, thì việc biểu diễn toán học của α và β có thể được dẫn xuất như sau

$$\alpha = \frac{I \cdot \sum_{i=1}^I (P_c[\mathbf{x}_i] \cdot P_r[\mathbf{x}_i]) - \sum_{i=1}^I (P_c[\mathbf{x}_i]) \cdot \sum_{i=1}^I (P_r[\mathbf{x}_i])}{I \cdot \sum_{i=1}^I (P_r[\mathbf{x}_i] \cdot P_r[\mathbf{x}_i]) - \left(\sum_{i=1}^I P_r[\mathbf{x}_i] \right)^2} \quad (12)$$

$$\beta = \frac{\sum_{i=1}^I (P_c[\mathbf{x}_i]) - \alpha \cdot \sum_{i=1}^I (P_r[\mathbf{x}_i])}{I}$$

Mà tại đó I biểu diễn số lượng các mẫu trong khuôn mẫu. $P_c[\mathbf{x}_i]$ là mẫu thứ i của khuôn mẫu của khối hiện tại và $P_r[\mathbf{x}_i]$ là mẫu tham chiếu của mẫu khuôn mẫu thứ i dựa trên véc-tơ chuyển động $\mathbf{v}$.

Bên cạnh việc được áp dụng cho các khối liên ảnh đều vốn tại chứa nhiều nhất một vectơ chuyển động cho mỗi hướng dự đoán (L0 hoặc L1), thì LIC cũng được áp dụng cho các khối được tạo mã chế độ afin mà tại đó một khối tạo mã còn được chia tách thành nhiều khối con nhỏ hơn và mỗi khối con có thể được kết hợp với thông tin chuyển động khác nhau. Để dẫn xuất các mẫu tham chiếu cho LIC của khối được tạo mã chế độ afin, như được thể hiện trên Fig.17A và Fig.17B (được mô tả dưới đây), thì các mẫu tham chiếu trong khuôn mẫu trên cùng của một khối tạo mã afin được tìm nạp nhờ sử dụng vectơ chuyển động của mỗi khối con trong hàng khối con trên cùng trong khi các mẫu tham chiếu trong khuôn mẫu bên trái được tìm nạp nhờ sử dụng các vectơ chuyển động của các khối con trong cột khối con bên trái. Sau việc đó, thì cùng một phương pháp dẫn xuất LLMSE như được thể hiện trong (12) sẽ được áp dụng để dẫn xuất các thông số LIC dựa trên khuôn mẫu tổng hợp.

Fig.17A thể hiện sự minh họa để dẫn xuất các mẫu khuôn mẫu cho chế độ afin, theo sáng chế. Sự minh họa này chứa Cur Frame 1720 và Cur CU 1722. Cur Frame 1720 là khung hiện tại. Cur CU 1722 là đơn vị tạo mã hiện tại.

Fig.17B thể hiện sự minh họa để dẫn xuất các mẫu khuôn mẫu cho chế độ afin. Sự minh họa này chứa Ref Frame 1740, Col CU 1742, A Ref 1743, B Ref 1744, C Ref 1745, D Ref 1746, E Ref 1747, F Ref 1748, và G Ref 1749. Ref Frame 1740 là khung tham chiếu. Col CU 1742 là một đơn vị tạo mã được đặt vào chỗ. A Ref 1743, B Ref 1744, C Ref 1745, D Ref 1746, E Ref 1747, F Ref 1748, và G Ref 1749 là các mẫu tham chiếu. Tính kém hiệu quả của sự tinh chỉnh dự đoán với luồng quang học cho chế độ afin

Mặc dù PROF có thể tăng cường hiệu quả tạo mã của chế độ afin, nhưng thiết kế của nó vẫn có thể được cải thiện thêm nữa. Đặc biệt là, với thực tế là PROF và BDOF đều được xây dựng dựa vào khái niệm luồng quang học, nên rất mong muốn để làm hài hòa

các thiết kế của PROF và BDOF càng nhiều càng tốt sao cho PROF có thể tận dụng tối đa các logic sẵn có của BDOF để tạo thuận lợi cho các cách triển khai phần cứng. Dựa trên sự xem xét như vậy, thì các tính kém hiệu quả về sự tương tác giữa các thiết kế PROF và BDOF hiện tại sẽ được nhận dạng trong sáng chế này.

Trước tiên, như được mô tả trong phần “sự tinh chỉnh dự đoán với luồng quang học cho chế độ affin”, thì trong phương trình (8), độ chính xác là các gradient được xác định dựa trên độ sâu bit bên trong. Mặt khác, sự chênh lệch MV, tức là, Δv_x và Δv_y , luôn được dẫn xuất tại độ chính xác là 1/32-pel. Theo cách tương ứng, dựa trên phương trình (9), thì độ chính xác của sự tinh chỉnh PROF được dẫn xuất phụ thuộc vào độ sâu bit bên trong. Tuy nhiên, tương tự với BDOF, PROF được áp dụng trên phần trên cùng của các giá trị mẫu dự đoán tại độ sâu bit cao trung gian (tức là, 16-bit) để giữ độ chính xác dẫn xuất PROF. Do đó, bất kể độ sâu bit tạo mã bên trong, thì độ chính xác của các sự tinh chỉnh dự đoán được dẫn xuất bởi PROF nên khớp với độ chính xác của các mẫu dự đoán trung gian, tức là, 16-bit. Nói cách khác, các độ sâu bit biểu diễn của sự chênh lệch MV và các gradient trong thiết kế PROF sẵn có không được khớp một cách hoàn hảo để dẫn xuất các sự tinh chỉnh dự đoán chuẩn xác tương đối với độ chính xác mẫu dự đoán (tức là, 16-bit). Trong lúc đó, dựa trên việc so sánh của các phương trình (1), (4) và (8), thì PROF và BDOF sẵn có sử dụng các độ chính xác khác nhau để biểu diễn các gradient mẫu và sự chênh lệch MV. Như được nêu ra trước đây, thiết kế không thống nhất như vậy sẽ không được mong muốn cho phần cứng bởi vì logic BDOF sẵn có không thể được sử dụng lại.

Thứ hai, như được thảo luận trong phần “sự tinh chỉnh dự đoán với luồng quang học cho chế độ affin”, khi một khối affin hiện tại được dự đoán đôi, thì PROF được áp dụng cho các mẫu dự đoán trong danh sách L0 và L1 theo kiểu tách biệt; sau đó, các tín

hiệu dự đoán L0 và L1 được tăng cường được tính trung bình để tạo ra tín hiệu dự đoán đôi cuối cùng. Ngược lại, thay vì dẫn xuất theo kiểu tách biệt sự tinh chỉnh PROF cho mỗi hướng dự đoán, thì BDOF dẫn xuất sự tinh chỉnh dự đoán một lần mà sau đó được áp dụng để tăng cường tín hiệu dự đoán L0 và L1 được tổ hợp. Fig.8 và Fig.9 (được mô tả dưới đây) so sánh tiến trình công việc của BDOF hiện tại và PROF cho việc dự đoán đôi. Trong thiết kế thiết kế đường ống phần cứng bộ mã hóa-giải mã thực tế, thì người ta thường chỉ định các môđun mã hóa/giải mã chính khác nhau cho mỗi giai đoạn đường ống sao cho nhiều khối tạo mã có thể được xử lý song song. Tuy nhiên, do sự chênh lệch giữa các tiến trình công việc BDOF và PROF, nên việc này có thể dẫn đến việc sẽ khó khăn để có cùng một thiết kế đường ống mà có thể được chia sẻ bởi BDOF và PROF, vốn không thân thiện với cách triển khai bộ mã hóa-giải mã thực tế.

Fig.8 thể hiện tiến trình công việc của BDOF, theo sáng chế. Tiến trình công việc 800 gồm việc bù chuyển động L0 810, việc bù chuyển động L1 820, và BDOF 830. Việc bù chuyển động L0 810, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu trước đó. Ảnh tham chiếu trước đó là ảnh tham chiếu trước đó từ ảnh hiện tại trong khối video. Việc bù chuyển động L1 820, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu tiếp theo. Ảnh tham chiếu tiếp theo là ảnh tham chiếu sau ảnh hiện tại trong khối video. BDOF 830 lấy các mẫu bù chuyển động từ việc bù chuyển động L1 810 và việc bù chuyển động L1 820 và xuất ra các mẫu dự đoán, như được mô tả đối với Fig.4 ở trên.

Fig.9 thể hiện tiến trình công việc của PROF sẵn có, theo sáng chế. Tiến trình công việc 900 gồm việc bù chuyển động L0 910, việc bù chuyển động L1 920, L0 PROF 930, L1 PROF 940, và trung bình 960. Việc bù chuyển động L0 910, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu trước đó. Ảnh tham chiếu trước đó

là ảnh tham chiếu trước đó từ ảnh hiện tại trong khối video. Việc bù chuyển động L1 920, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu tiếp theo. Ảnh tham chiếu tiếp theo là ảnh tham chiếu sau ảnh hiện tại trong khối video. L0 PROF 930 lấy các mẫu bù chuyển động L0 từ việc bù chuyển động L0 910 và xuất ra các giá trị tinh chỉnh chuyển động, như được mô tả đối với Fig.7 ở trên. L1 PROF 940 lấy việc bù chuyển động L1 các mẫu từ việc bù chuyển động L1 920 và xuất ra các giá trị tinh chỉnh chuyển động, như được mô tả đối với Fig.7 ở trên. Trung bình 960 tính trung bình các đầu ra giá trị tinh chỉnh chuyển động của L0 PROF 930 và L1 PROF 940.

Thứ ba, đối với cả BDOF và PROF, thì các gradient cần được tính toán cho mỗi mẫu bên trong khối tạo mã hiện tại, vốn đòi hỏi tạo ra một hàng/cột bổ sung của các mẫu dự đoán trên mỗi phía của khối. Để tránh độ phức tạp tính toán bổ sung của phép nội suy mẫu, thì các mẫu dự đoán trong vùng được mở rộng quanh khối được sao chép một cách trực tiếp từ các mẫu tham chiếu tại vị trí nguyên (tức là, không cần phép nội suy). Tuy nhiên, theo thiết kế sẵn có, thì các mẫu nguyên tại các sự định vị khác nhau được lựa chọn để tạo ra các giá trị gradient của BDOF và PROF. Cụ thể là, đối với BDOF, thì mẫu tham chiếu nguyên mà được đặt bên trái của mẫu dự đoán (đối với các gradient ngang) và ở trên mẫu dự đoán (đối với các gradient dọc) được sử dụng; đối với PROF, thì mẫu tham chiếu nguyên mà gần nhất với mẫu dự đoán được sử dụng cho các sự tính toán gradient. Tương tự với vấn đề biểu diễn độ sâu bit, thì phương pháp tính toán gradient không thống nhất này cũng là không mong muốn đối với các cách triển khai bộ mã hóa-giải mã.

Thứ tư, như được nêu ra trước đây, động lực của PROF là để bù cho sự chênh lệch MV nhỏ giữa MV của mỗi mẫu và MV khối con mà được dẫn xuất tại trung tâm của khối con mà mẫu thuộc về. Theo thiết kế PROF hiện tại, thì PROF luôn được gọi khi

một khối tạo mã được dự đoán bởi chế độ affin. Tuy nhiên, như được chỉ ra trong phương trình (6) và (7), các MV khối con của một khối affin được dẫn xuất từ các MV điểm điều khiển. Do đó, khi sự chênh lệch giữa các MV điểm điều khiển là tương đối nhỏ, thì các MV tại mỗi vị trí mẫu là phải nhất quán. Trong trường hợp như vậy, bởi vì lợi ích của áp dụng PROF có thể rất bị hạn chế, nên việc thực hiện PROF có thể không đáng giá khi xem xét đến sự cân bằng hiệu suất/độ phức tạp.

Các sự cải thiện đối với sự tinh chỉnh dự đoán với luồng quang học cho chế độ affin

Trong sáng chế này, các phương pháp được đề xuất để cải thiện và đơn giản hóa thiết kế PROF sẵn có để tạo thuận lợi các cách triển khai bộ mã hóa-giải mã phân cứng. Một cách cụ thể, cần đặc biệt chú ý đến việc làm hài hòa các thiết kế của BDOF và PROF để chia sẻ một cách tối đa các logic BDOF sẵn có với PROF. Nói chung, các khía cạnh chính của các công nghệ được đề xuất trong sáng chế này được tóm tắt như sau.

Thứ nhất, để cải thiện hiệu quả tạo mã của PROF trong khi đạt được một thiết kế thống nhất hơn, thì một phương pháp được đề xuất để thống nhất độ sâu bit biểu diễn của các gradient mẫu và sự chênh lệch MV mà được sử dụng bởi BDOF và PROF.

Thứ hai, để tạo thuận lợi cho thiết kế đường ống phân cứng, thì sáng chế đề xuất làm hài hòa tiến trình công việc của PROF với tiến trình công việc của BDOF cho việc dự đoán đôi. Cụ thể là, không giống với PROF sẵn có mà dẫn xuất các sự tinh chỉnh dự đoán theo kiểu tách biệt cho L0 và L1, thì phương pháp dẫn xuất sự tinh chỉnh dự đoán được đề xuất một lần mà được áp dụng cho tín hiệu dự đoán L0 và L1 được tổ hợp.

Thứ ba, hai phương pháp được đề xuất để làm hài hòa việc dẫn xuất của các mẫu tham chiếu nguyên để tính toán các giá trị gradient mà được sử dụng bởi BDOF và PROF.

Thứ tư, để giảm độ phức tạp tính toán, thì các phương pháp kết thúc sớm được đề xuất để không cho phép theo cách thích ứng quy trình PROF đối với các khối tạo mã afin khi các điều kiện nhất định được thỏa mãn.

Thiết kế biểu diễn độ sâu bit được cải thiện của các gradient PROF và sự chênh lệch MV

Như được phân tích trong phần “báo cáo vấn đề”, các độ sâu bit biểu diễn của sự chênh lệch MV và các gradient mẫu trong PROF hiện tại không được căn chỉnh để dẫn xuất các sự tinh chỉnh dự đoán chuẩn xác. Hơn thế nữa, độ sâu bit biểu diễn của các gradient mẫu và sự chênh lệch MV là không nhất quán giữa BDOF và PROF, vốn không thân thiện đối với phần cứng. Trong phần này, một phương pháp biểu diễn độ sâu bit được cải thiện sẽ được đề xuất bởi phương pháp biểu diễn độ sâu bit mở rộng của BDOF thành PROF. Cụ thể là, trong phương pháp được đề xuất, các gradient ngang và dọc tại mỗi vị trí mẫu được tính toán dưới dạng

$$\begin{aligned} g_x(i, j) &= (I(i + 1, j) - I(i - 1, j)) \gg \max(6, \text{bitdepth} - 6) \\ g_y(i, j) &= (I(i, j + 1) - I(i, j - 1)) \gg \max(6, \text{bitdepth} - 6) \end{aligned} \quad (13)$$

Theo cách bổ sung, giả sử Δx và Δy là dịch vị ngang và dọc được biểu diễn tại độ chuẩn xác $\frac{1}{4}$ -pel từ một sự định vị mẫu đến trung tâm của khối con mà mẫu thuộc về, thì sự chênh lệch MV PROF tương ứng $\Delta v(x, y)$ tại vị trí mẫu được dẫn xuất dưới dạng

$$\begin{aligned} \Delta v_x(i, j) &= (c * \Delta x + d * \Delta y) \gg (13 - dMvBits) \\ \Delta v_y(i, j) &= (e * \Delta x + f * \Delta y) \gg (13 - dMvBits) \end{aligned} \quad (14)$$

mà tại đó $dMvBits$ là độ sâu bit của các giá trị gradient mà được sử dụng bởi quy trình BDOF, tức là, $dMvBits = \max(5, (\text{bitdepth} - 7)) + 1$. Trong phương trình (13) và (14), c , d , e và f là các thông số afin vốn được dẫn xuất dựa trên các MV điểm điều khiển afin. Cụ thể là, đối với mô hình afin 4-thông số,

$$\begin{cases} c = f = \frac{v_{1x} - v_{0x}}{w} \\ e = -d = \frac{v_{1y} - v_{0y}}{w} \end{cases}$$

Đối với mô hình afin 6-thông số,

$$\begin{cases} c = \frac{v_{1x} - v_{0x}}{w} \\ d = \frac{v_{2x} - v_{0x}}{h} \\ e = \frac{v_{1y} - v_{0y}}{w} \\ f = \frac{v_{2y} - v_{0y}}{h} \end{cases}$$

mà tại đó (v_{0x}, v_{0y}) , (v_{1x}, v_{1y}) , (v_{2x}, v_{2y}) là các MV điểm điều khiển trên cùng bên trái, trên cùng bên phải và dưới cùng bên trái của khối tạo mã hiện tại vốn được biểu diễn trong độ chính xác 1/16-pel, và w và h là chiều rộng và chiều cao của khối.

Trong phần thảo luận ở trên, như được thể hiện trong phương trình (13) và (14), một cặp các dịch chuyển sang phải cố định được áp dụng để tính toán các giá trị của các gradient và các sự chênh lệch MV. Trong thực tế, các dịch chuyển sang phải đảo bit khác nhau có thể được áp dụng cho (13) và (14) để đạt được các sự chính xác biểu diễn khác nhau của các gradient và sự chênh lệch MV cho sự cân bằng khác nhau giữa độ chính xác tính toán trung gian và chiều rộng bit của quy trình dẫn xuất PROF bên trong. Ví dụ, khi video đầu vào chứa nhiều nhiễu, thì các gradient được dẫn xuất có thể không đáng tin cậy để biểu diễn các giá trị ngang/dọc gradient cục bộ đúng tại mỗi mẫu. Trong trường hợp như vậy, sẽ hợp lý outputs khi sử dụng nhiều bit hơn để biểu diễn sự chênh lệch MV so với các gradient. Mặt khác, khi video đầu vào thể hiện chuyển động ổn định, thì các sự chênh lệch MV như được dẫn xuất bởi mô hình afin sẽ rất nhỏ. Nếu vậy, việc sử dụng sự chênh lệch MV chính xác cao không thể cung cấp lợi ích bổ sung để tăng độ

chính xác của sự tinh chỉnh PROF được dẫn xuất. Nói cách khác, trong trường hợp như vậy, sẽ có lợi hơn khi sử dụng nhiều bit hơn để biểu diễn các giá trị gradient. Dựa trên phân xem xét ở trên, trong một hoặc nhiều phương án của sáng chế, thì một phương pháp chung như được đề xuất trong phần sau đây để tính toán các gradient và sự chênh lệch MV đối với PROF. Cụ thể là, giả sử các gradient ngang và dọc tại mỗi vị trí mẫu được tính toán bằng cách áp dụng các dịch chuyển sang phải n_a cho sự chênh lệch của các mẫu dự đoán lân cận, tức là,

$$\begin{aligned} g_x(i, j) &= (I(i + 1, j) - I(i - 1, j)) \gg n_a \\ g_y(i, j) &= (I(i, j + 1) - I(i, j - 1)) \gg n_a \end{aligned} \quad (15)$$

thì sự chênh lệch MV PROF tương ứng $\Delta v(x, y)$ tại vị trí mẫu nên được tính toán dưới dạng

$$\begin{aligned} \Delta v_x(i, j) &= (c * \Delta x + d * \Delta y) \gg (13 - n_a) \\ \Delta v_y(i, j) &= (e * \Delta x + f * \Delta y) \gg (13 - n_a) \end{aligned} \quad (16)$$

mà tại đó Δx và Δy là dịch vị ngang và dọc được biểu diễn tại độ chuẩn xác $1/4$ -pel từ một sự định vị mẫu đến trung tâm của khối con mà mẫu thuộc về và c, d, e và f là các thông số afin vốn được dẫn xuất dựa trên các MV điểm điều khiển afin $1/16$ -pel. Cuối cùng, sự tinh chỉnh PROF cuối cùng của mẫu được tính toán dưới dạng

$$\Delta I(i, j) = (g_x(i, j) * \Delta v_x(i, j) + g_y(i, j) * \Delta v_y(i, j) + 1) \gg 1 \quad (17)$$

Trong một phương án khác, một phương pháp điều khiển độ sâu bit PROF khác được đề xuất như sau. Trong phương pháp, các gradient ngang và dọc tại mỗi vị trí mẫu vẫn được tính toán như trong (18) bằng cách áp dụng n_a bit của các dịch chuyển sang phải cho giá trị chênh lệch của các mẫu dự đoán lân cận. Sự chênh lệch MV PROF tương ứng $\Delta v(x, y)$ tại vị trí mẫu nên được tính toán dưới dạng:

$$\begin{aligned}\Delta v_x(i,j) &= (c * \Delta x + d * \Delta y) \gg (14 - n_a) \\ \Delta v_y(i,j) &= (e * \Delta x + f * \Delta y) \gg (14 - n_a)\end{aligned}$$

Theo cách bổ sung, để giữ toàn bộ việc dẫn xuất PROF tại độ sâu bit bên trong thích hợp, thì việc cắt xén được áp dụng cho sự chênh lệch MV được dẫn xuất như sau:

$$\begin{aligned}\Delta v_x(i,j) &= \text{Clip3}(-limit, limit, \Delta v_x(i,j)) \\ \Delta v_y(i,j) &= \text{Clip3}(-limit, limit, \Delta v_y(i,j))\end{aligned}$$

mà tại đó giới hạn là ngưỡng mà bằng 2^{n_b} và $\text{clip3}(min, max, x)$ là hàm mà cắt xén giá trị x đã cho bên trong phạm vi $[min, max]$. Trong một ví dụ, giá trị của n_b được thiết lập là $2^{\max(5, bitdepth - 7)}$. Cuối cùng, sự tinh chỉnh PROF của mẫu được tính toán dưới dạng

$$\Delta I(i,j) = g_x(i,j) * \Delta v_x(i,j) + g_y(i,j) * \Delta v_y(i,j)$$

Theo cách bổ sung, trong một hoặc nhiều phương án, sáng chế đề xuất một giải pháp điều khiển độ sâu bit PROF. Trong phương pháp này, các sự tinh chỉnh chuyển động PROF ngang và dọc tại mỗi vị trí mẫu (i, j) được dẫn xuất dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= (c * \Delta x + d * \Delta y) \gg (13 - \max(5, bit\ depth - 7)) \\ \Delta v_y(i,j) &= (e * \Delta x + f * \Delta y) \gg (13 - \max(5, bit\ depth - 7))\end{aligned}$$

Hơn nữa, các sự tinh chỉnh chuyển động ngang và dọc được dẫn xuất sẽ được cắt xén dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= \text{Clip3}(-\max(5, bit\ depth - 7), \max(5, bit\ depth - 7) - 1, \Delta v_x(i,j)) \\ \Delta v_y(i,j) &= \text{Clip3}(-\max(5, bit\ depth - 7), \max(5, bit\ depth - 7) - 1, \Delta v_y(i,j))\end{aligned}$$

Ở đây, căn cứ vào các sự tinh chỉnh chuyển động như được dẫn xuất ở trên, thì sự tinh chỉnh mẫu PROF cuối cùng tại sự định vị (i, j) được tính toán dưới dạng

$$\Delta I(i,j) = g_x(i,j) * \Delta v_x(i,j) + g_y(i,j) * \Delta v_y(i,j)$$

Trong một phương án khác, sáng chế đề xuất một giải pháp điều khiển độ sâu bit PROF khác. Trong phương pháp thứ hai, các sự tinh chỉnh chuyển động PROF ngang và dọc tại vị trí mẫu (i, j) được dẫn xuất dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= (c * \Delta x + d * \Delta y) \gg (13 - \max(6, \text{bit depth} - 6)) \\ \Delta v_y(i,j) &= (e * \Delta x + f * \Delta y) \gg (13 - \max(6, \text{bit depth} - 6))\end{aligned}$$

Sau đó, các sự tinh chỉnh chuyển động được dẫn xuất sẽ được cắt xén dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= \text{Clip3}(-\max(5, \text{bit depth} - 7), \max(5, \text{bit depth} - 7) - 1, \Delta v_x(i,j)) \\ \Delta v_y(i,j) &= \text{Clip3}(-\max(5, \text{bit depth} - 7), \max(5, \text{bit depth} - 7) - 1, \Delta v_y(i,j))\end{aligned}$$

Vì thế, căn cứ vào các sự tinh chỉnh chuyển động như được dẫn xuất ở trên, thì sự tinh chỉnh mẫu PROF cuối cùng tại sự định vị (i, j) được tính toán dưới dạng

$$\Delta I(i,j) = (g_x(i,j) * \Delta v_x(i,j) + g_y(i,j) * \Delta v_y(i,j) + 1) \gg 1$$

Trong một hoặc nhiều phương án, sáng chế đề xuất tổ hợp phương pháp điều khiển chính xác tinh chỉnh chuyển động trong giải pháp và phương pháp dẫn xuất tinh chỉnh mẫu PROF trong giải pháp thứ hai. Cụ thể là, theo phương pháp này, thì sự tinh chỉnh chuyển động PROF ngang và dọc tại mỗi vị trí mẫu (i, j) được dẫn xuất dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= (c * \Delta x + d * \Delta y) \gg (13 - \max(5, \text{bit depth} - 7)) \\ \Delta v_y(i,j) &= (e * \Delta x + f * \Delta y) \gg (13 - \max(5, \text{bit depth} - 7))\end{aligned}$$

Hơn nữa, các sự tinh chỉnh chuyển động ngang và dọc được dẫn xuất được cắt xén dưới dạng

$$\begin{aligned}\Delta v_x(i,j) &= \text{Clip3}(-\max(5, \text{bit depth} - 7), \max(5, \text{bit depth} - 7) - 1, \Delta v_x(i,j)) \\ \Delta v_y(i,j) &= \text{Clip3}(-\max(5, \text{bit depth} - 7), \max(5, \text{bit depth} - 7) - 1, \Delta v_y(i,j))\end{aligned}$$

Ở đây, căn cứ vào các sự tinh chỉnh chuyển động như được dẫn xuất ở trên, thì sự tinh chỉnh mẫu PROF cuối cùng tại sự định vị (i, j) được tính toán dưới dạng

$$\Delta I(i,j) = (g_x(i,j) * \Delta v_x(i,j) + g_y(i,j) * \Delta v_y(i,j) + 1) \gg 1$$

Trong một hoặc nhiều phương án, thì phương pháp dẫn xuất tinh chỉnh mẫu PROF sau đây được đề xuất:

Thứ nhất, tính toán các sự tinh chỉnh chuyển động ngang và dọc PROF nằm trong độ chính xác là 1/32-pel bằng cách áp dụng các dịch chuyển sang phải cố định sau đây, như được chỉ ra dưới dạng

$$\begin{aligned}\Delta v_x(i, j) &= (c * \Delta x + d * \Delta y) \gg 8 \\ \Delta v_y(i, j) &= (e * \Delta x + f * \Delta y) \gg 8\end{aligned}$$

Thứ hai, cắt xén các giá trị tinh chỉnh chuyển động PROF được tính toán thành một phạm vi đối xứng $[-31, 31]$.

$$\begin{aligned}\Delta v_x(i, j) &= \text{Clip3}(-31, 31, \Delta v_x(i, j)) \\ \Delta v_y(i, j) &= \text{Clip3}(-31, 31, \Delta v_y(i, j))\end{aligned}$$

Thứ ba, sự tinh chỉnh PROF của mẫu được tính toán dưới dạng

$$\Delta I(i, j) = g_x(i, j) * \Delta v_x(i, j) + g_y(i, j) * \Delta v_y(i, j)$$

Fig.10 thể hiện phương pháp biểu diễn độ sâu bit của PROF. Phương pháp có thể, ví dụ, được áp dụng cho bộ giải mã.

Trong bước 1010, bộ giải mã có thể thu nhận ảnh tham chiếu I được kết hợp với khối video trong tín hiệu video.

Trong bước 1012, bộ giải mã có thể thu nhận các mẫu dự đoán $I(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu I . i và j có thể biểu diễn tọa độ của một mẫu với khối video.

Trong bước 1014, bộ giải mã có thể điều khiển các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng việc dịch chuyển sang phải cho các thông số PROF bên trong dựa trên giá trị dịch chuyển bit để đạt được sự chính xác được thiết

lập trước. Các thông số PROF bên trong gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu dự đoán $I(i, j)$.

Trong bước 1016, bộ giải mã có thể thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên quy trình dẫn xuất PROF đang được áp dụng cho khối video dựa trên các mẫu dự đoán $I(i, j)$.

Trong bước 1018, bộ giải mã có thể thu nhận các mẫu dự đoán của khối video dựa trên tổ hợp của các mẫu dự đoán và các giá trị tinh chỉnh dự đoán.

Theo cách bổ sung, cùng một phương pháp dẫn xuất thông số cũng có thể được áp dụng cho quy trình tinh chỉnh mẫu BDOF như được minh họa dưới dạng

$$v_x = sGx2 > 0 ? \text{Clip3}(-31, 31,$$

$$-(sGxdI \ll 2) \gg \text{Floor}(\text{Log2}(sGx2))) : 0$$

$$v_y = sGy2 > 0 ? \text{Clip3}(-31, 31, ((sGydI \ll 2) -$$

$$((v_x * sGxGym) \ll 12 + v_x * sGxGys) \gg 1) \gg \text{Floor}(\text{Log2}(sGy2))) : 0$$

mà tại đó $sGxdI$, $sGx2$, $sGxGym$, $sGxGys$ và $sGy2$ là các thông số dẫn xuất BDOF trung gian.

Fig.11 thể hiện phương pháp điều khiển độ sâu bit của BDOF. Phương pháp có thể, ví dụ, được áp dụng cho bộ giải mã.

Trong bước 1110, bộ giải mã có thể thu nhận ảnh tham chiếu thứ nhất $I^{(0)}$ và ảnh tham chiếu thứ hai $I^{(1)}$ được kết hợp với khối video. Ảnh tham chiếu thứ nhất $I^{(0)}$ ở trước ảnh hiện tại và ảnh tham chiếu thứ hai $I^{(1)}$ ở sau ảnh hiện tại trong thứ tự hiển thị.

Trong bước 1112, bộ giải mã có thể thu nhận các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ nhất $I^{(0)}$. i và j có thể biểu diễn tọa độ của một mẫu với ảnh hiện tại.

Trong bước 1114, bộ giải mã có thể thu nhận các mẫu dự đoán thứ hai $I^{(1)}$ của khối video từ khối tham chiếu ở trong ảnh tham chiếu thứ hai $I^{(1)}$.

Trong bước 1116, bộ giải mã có thể điều khiển các thông số BDOF bên trong của quy trình dẫn xuất BDOF bằng cách áp dụng việc dịch chuyển cho các thông số BDOF bên trong. Các thông số BDOF bên trong bao gồm các giá trị gradient ngang và các giá trị gradient dọc được dẫn xuất dựa trên các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$, các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, các sự chênh lệch mẫu giữa các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ và các mẫu dự đoán thứ hai $I^{(1)}(i, j)$, và các thông số dẫn xuất BDOF trung gian. Các thông số dẫn xuất BDOF trung gian gồm các thông số sGxdI, sGydI, sGx2, sGxGy, và sGy2. sGxdI và sGydI gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị chênh lệch mẫu, và giữa các giá trị gradient dọc và các giá trị chênh lệch mẫu. sGx2 và sGy2 gồm các giá trị tương liên tự động của các giá trị gradient ngang và các giá trị gradient dọc. sGxGy gồm các giá trị tương liên chéo giữa các giá trị gradient ngang và các giá trị gradient dọc.

Trong bước 1118, bộ giải mã có thể thu nhận các sự tinh chỉnh chuyển động cho các mẫu trong khối video dựa trên BDOF đang được áp dụng cho khối video dựa trên các mẫu dự đoán thứ nhất $I^{(0)}(i, j)$ và các mẫu dự đoán thứ hai $I^{(1)}(i, j)$.

Trong bước 1120, bộ giải mã có thể thu nhận các mẫu dự đoán đôi của khối video dựa trên các sự tinh chỉnh chuyển động.

Các tiến trình công việc được làm hài hòa của BDOF và PROF cho việc dự đoán đôi

Như được thảo luận trước đây, khi một khối tạo mã afin được dự đoán đôi, thì PROF hiện tại được áp dụng theo cách thức đơn phương. Cụ thể hơn là, các sự tinh chỉnh mẫu PROF được dẫn xuất và được áp dụng theo kiểu tách biệt cho các mẫu dự đoán trong danh sách L0 và L1. Sau việc đó, thì các tín hiệu dự đoán được tinh chỉnh, một cách tương ứng từ danh sách L0 và L1, được tính trung bình để tạo ra tín hiệu dự đoán đôi cuối cùng của khối. Điều này là trái ngược với thiết kế BDOF mà tại đó các sự tinh chỉnh mẫu được dẫn xuất và được áp dụng cho tín hiệu dự đoán đôi. Sự chênh lệch như vậy giữa các tiến trình công việc dự đoán đôi của BDOF và PROF có thể không thân thiện với thiết kế đường ống bộ mã hóa-giải mã thực tế.

Để tạo thuận lợi cho thiết kế đường ống phần cứng, thì một phương pháp đơn giản hóa theo sáng chế theo là để sửa đổi quy trình dự đoán đôi của PROF sao cho các tiến trình công việc của hai phương pháp tinh chỉnh dự đoán này được làm hài hòa. Cụ thể là, thay vì áp dụng theo kiểu tách biệt sự tinh chỉnh cho mỗi hướng dự đoán, thì phương pháp PROF được đề xuất dẫn xuất các sự tinh chỉnh dự đoán một lần dựa trên các MV điểm điều khiển của danh sách L0 và L1; sau đó các sự tinh chỉnh dự đoán được dẫn xuất sẽ được áp dụng cho tín hiệu dự đoán L0 và L1 được tổ hợp để tăng cường chất lượng. Cụ thể là, dựa trên sự chênh lệch MV như được dẫn xuất trong phương trình (14), thì các mẫu dự đoán đôi cuối cùng của một khối tạo mã afin được tính toán bởi phương pháp được đề xuất dưới dạng

$$\begin{aligned} pred_{PROF}(i,j) &= (I^{(0)}(i,j) + I^{(1)}(i,j) + \Delta I(i,j) + o_{offset}) \gg shift \\ \Delta I(i,j) &= (g_x(i,j) * \Delta v_x(i,j) + g_y(i,j) * \Delta v_y(i,j) + 1) \gg 1 \\ I^r(i,j) &= I(i,j) + \Delta I(i,j) \end{aligned} \quad (18)$$

mà tại đó $shift$ và o_{offset} là giá trị dịch chuyển sang phải và giá trị dịch vị mà được áp dụng để tổ hợp các tín hiệu dự đoán L0 và L1 cho việc dự đoán đôi, vốn bằng với $(15 - bitdepth)$ và $1 \ll (14 - bitdepth) + (2 \ll 13)$, một cách tương ứng. Hơn thế nữa,

như được thể hiện trong (18), phép toán cắt xén trong thiết kế PROF sẵn có (như được thể hiện trong (9)) được xóa bỏ trong phương pháp được đề xuất.

Fig.13 minh họa quy trình PROF tương ứng khi phương pháp PROF dự đoán đôi được đề xuất sẽ được áp dụng. Quy trình PROF 1300 gồm việc bù chuyển động L0 1310, việc bù chuyển động L1 1320, và PROF việc dự đoán đôi 1330. Việc bù chuyển động L0 1310, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu trước đó. Ảnh tham chiếu trước đó là ảnh tham chiếu trước đó từ ảnh hiện tại trong khối video. Việc bù chuyển động L1 1320, ví dụ, có thể là danh sách của các mẫu bù chuyển động từ ảnh tham chiếu tiếp theo. Ảnh tham chiếu tiếp theo là ảnh tham chiếu sau ảnh hiện tại trong khối video. PROF việc dự đoán đôi 1330 lấy các mẫu bù chuyển động từ việc bù chuyển động L1 1310 và việc bù chuyển động L1 1320 và xuất ra các mẫu dự đoán đôi, như được mô tả ở trên.

Để chứng minh lợi ích tiềm năng của phương pháp được đề xuất cho việc thiết kế đường ống phân cứng, thì Fig.14 (được mô tả dưới đây) thể hiện một ví dụ để minh họa giai đoạn đường ống khi cả BDOF và PROF được đề xuất đều được áp dụng. Trên Fig.14, quy trình giải mã của một khối liên ảnh chủ yếu chứa ba bước:

Thứ nhất, phân tích cú pháp/giải mã các MV của khối tạo mã và tìm nạp các mẫu tham chiếu.

Thứ hai, tạo ra các tín hiệu dự đoán L0 và/hoặc L1 của khối tạo mã.

Thứ ba, thực hiện sự tinh chỉnh đảo mẫu của các mẫu dự đoán đôi được tạo ra dựa trên BDOF khi khối tạo mã được dự đoán bởi một chế độ không-afin hoặc PROF khi khối tạo mã được dự đoán bởi chế độ affin.

Fig.14 thể hiện sự minh họa của giai đoạn đường ống ví dụ khi cả BDOF và PROF được đề xuất đều được áp dụng, theo sáng chế. Fig.14 chứng minh lợi ích tiềm năng của phương pháp được đề xuất cho việc thiết kế đường ống phần cứng. Giai đoạn đường ống 1400 gồm bước phân tích cú pháp/giải mã MV và tìm nạp các mẫu tham chiếu 1410, việc bù chuyển động 1420, BDOF/PROF 1430. Giai đoạn đường ống 1400 sẽ mã hóa các khối video BLK0, BKL1, BKL2, BKL3, và BLK4. Mỗi khối video sẽ bắt đầu trong việc phân tích cú pháp/giải mã MV và tìm nạp các mẫu tham chiếu 1410 và di chuyển đến việc bù chuyển động 1420 và sau đó việc bù chuyển động 1420, BDOF/PROF 1430, một cách tuần tự. Điều này có nghĩa là BLK0 sẽ không bắt đầu trong quy trình giai đoạn đường ống 1400 cho đến khi BLK0 di chuyển lên trên việc bù chuyển động 1420. Tương tự đối với tất cả các giai đoạn và các khối video khi thời gian đi từ T0 đến T1, T2, T3, và T4.

Như được thể hiện trên Fig.14, sau khi phương pháp làm hài hòa được đề xuất mà được áp dụng, thì cả BDOF và PROF đều được áp dụng một cách trực tiếp cho các mẫu dự đoán đôi. Căn cứ vào việc BDOF và PROF được áp dụng cho các loại khối tạo mã khác nhau (tức là, BDOF được áp dụng cho các khối không-afin và PROF được áp dụng cho các khối affin), thì hai công cụ tạo mã không thể được gọi đồng thời. Do đó, các quy trình giải mã tương ứng của chúng có thể được tiến hành bằng cách chia sẻ cùng một giai đoạn đường ống. Điều này hiệu quả hơn so với thiết kế PROF sẵn có mà tại đó khó để chỉ định cùng một giai đoạn đường ống cho cả BDOF và PROF do tiến trình công việc của việc dự đoán đôi khác nhau của chúng.

Trong phần thảo luận ở trên, phương pháp được đề xuất chỉ xem xét đến việc làm hài hòa các tiến trình công việc của BDOF và PROF. Tuy nhiên, theo các thiết kế sẵn có, thì đơn vị vận hành cơ bản cho hai công cụ tạo mã cũng được thực hiện tại các kích

cỡ khác nhau. Ví dụ, đối với BDOF, một khối tạo mã được chia tách thành nhiều khối con với kích cỡ là $W_s \times H_s$, mà tại đó $W_s = \min(W, 16)$ và $H_s = \min(H, 16)$, mà tại đó W và H là chiều rộng và chiều cao của khối tạo mã. Các phép toán BDOF, như là việc tính toán gradient và việc dẫn xuất tinh chỉnh mẫu, được thực hiện một cách độc lập đối với mỗi khối con. Mặt khác, như được mô tả trước đây, khối tạo mã afin được chia thành các khối con 4×4 , với mỗi khối con được chỉ định một MV riêng lẻ được dẫn xuất dựa trên hoặc các mô hình afin 4-thông số hoặc 6-thông số. Bởi vì PROF chỉ được áp dụng cho khối afin, nên đơn vị vận hành cơ bản của nó là khối con 4×4 . Tương tự với vấn đề tiến trình công việc dự đoán đôi, thì việc sử dụng kích cỡ đơn vị vận hành khác nhau đối với PROF từ BDOF cũng không thân thiện đối với các cách triển khai phần cứng và gây khó khăn đối với BDOF và PROF để chia sẻ cùng một giai đoạn đường ống của toàn bộ quy trình giải mã. Để giải quyết vấn đề này, trong một hoặc nhiều phương án, sáng chế đề xuất căn chỉnh kích cỡ khối con của chế độ afin để giống với kích cỡ khối con của BDOF. Ví dụ, theo phương pháp được đề xuất, nếu một khối tạo mã được tạo mã bởi chế độ afin, thì nó sẽ được chia tách thành các khối con với kích cỡ là $W_s \times H_s$, mà tại đó $W_s = \min(W, 16)$ và $H_s = \min(H, 16)$, mà tại đó W và H là chiều rộng và chiều cao của khối tạo mã. Mỗi khối con được chỉ định một MV riêng lẻ và được xem là một đơn vị vận hành PROF độc lập. Điều đáng nói là đơn vị vận hành PROF độc lập đảm bảo rằng phép toán PROF trên phần trên cùng của nó được thực hiện mà không tham chiếu đến thông tin từ các đơn vị vận hành PROF lân cận. Ví dụ, sự chênh lệch MV PROF tại một vị trí mẫu được tính toán dưới dạng sự chênh lệch giữa MV tại vị trí mẫu và MV tại trung tâm của đơn vị vận hành PROF mà mẫu được đặt trong đó; các gradient được sử dụng bởi việc dẫn xuất PROF được tính toán bằng cách đếm các mẫu dọc theo mỗi đơn vị vận hành PROF. Các lợi ích đã được khẳng định của phương pháp được đề xuất chủ yếu gồm các khía cạnh sau: 1) kiến trúc đường ống được đơn giản hóa với kích cỡ cơ

bản đơn vị vận hành thống nhất cho cả việc bù chuyển động và sự tinh chỉnh BDOF/PROF; 2) việc sử dụng băng thông bộ nhớ được giảm do kích cỡ khối con được phóng to cho việc bù chuyển động afin; 3) độ phức tạp tính toán được giảm trên mỗi mẫu của phép nội suy mẫu phân số.

Cũng nên đề cập đến rằng bởi vì độ phức tạp tính toán được giảm (tức là, mục 3)) với phương pháp được đề xuất, nên ràng buộc bộ lọc nội suy 6-tap sẵn có cho các khối tạo mã afin có thể được xóa bỏ. Thay vào đó, phép nội suy 8-tap mặc định cho các khối tạo mã không-afin cũng được sử dụng cho afin các khối tạo mã. Độ phức tạp tính toán tổng thể trong trường hợp này vẫn có thể so sánh một cách tích cực so với thiết kế PROF sẵn có (mà dựa trên khối con 4x4 với bộ lọc nội suy 6-tap).

Việc làm hài hòa của việc dẫn xuất gradient đối với BDOF và PROF

Như được mô tả trước đây, cả BDOF và PROF đều tính toán gradient của mỗi mẫu bên trong khối tạo mã hiện tại, vốn truy cập vào một hàng/cột bổ sung của các mẫu dự đoán trên mỗi phía của khối. Để tránh độ phức tạp nội suy bổ sung, thì các mẫu dự đoán cần có trong vùng được mở rộng quanh ranh giới khối được sao chép một cách trực tiếp từ các mẫu tham chiếu nguyên. Tuy nhiên, như được nêu ra trong phần “báo cáo vấn đề”, các mẫu nguyên tại các sự định vị khác nhau được sử dụng để tính toán các giá trị gradient của BDOF và PROF.

Để đạt được một thiết kế đồng nhất hơn, thì hai phương pháp được đề xuất trong phần sau đây để thống nhất phương pháp dẫn xuất các gradient sẽ được sử dụng bởi BDOF và PROF. Trong phương pháp thứ nhất, sáng chế đề xuất căn chỉnh phương pháp dẫn xuất gradient của PROF để giống với phương pháp dẫn xuất gradient của BDOF. Ví dụ, theo phương pháp thứ nhất, vị trí nguyên được sử dụng để tạo ra các mẫu dự đoán trong vùng được mở rộng sẽ được xác định bằng cách đặt sàn xuống vị trí mẫu phân số,

tức là, vị trí mẫu nguyên được lựa chọn sẽ được đặt bên trái vị trí mẫu phân số (đối với các gradient ngang) và ở trên vị trí mẫu phân số (đối với các gradient dọc). Trong phương pháp thứ hai, sáng chế đề xuất căn chỉnh phương pháp dẫn xuất gradient của BDOF để giống với phương pháp dẫn xuất gradient của PROF. Chi tiết hơn là, khi phương pháp thứ hai được áp dụng, thì mẫu tham chiếu nguyên mà gần nhất với mẫu dự đoán sẽ được sử dụng cho các sự tính toán gradient.

Fig.15 thể hiện ví dụ về việc sử dụng phương pháp dẫn xuất gradient của BDOF, theo sáng chế. Trên Fig.15, các hình tròn trắng 1510 biểu diễn các mẫu tham chiếu tại các vị trí nguyên, các hình tam giác 1530 biểu diễn các mẫu dự đoán phân số của khối hiện tại, và các hình tròn màu đen 1520 biểu diễn các mẫu tham chiếu nguyên mà được sử dụng để lấp đầy vùng được mở rộng của khối hiện tại.

Fig.16 thể hiện ví dụ về việc sử dụng phương pháp dẫn xuất gradient của PROF, theo sáng chế. Fig.16, các hình tròn trắng 1610 biểu diễn các mẫu tham chiếu tại các vị trí nguyên, các hình tam giác 1630 biểu diễn các mẫu dự đoán phân số của khối hiện tại, và các hình tròn màu đen 1620 biểu diễn các mẫu tham chiếu nguyên mà được sử dụng để lấp đầy vùng được mở rộng của khối hiện tại.

Fig.15 và Fig.16 minh họa các sự định vị mẫu nguyên tương ứng mà được sử dụng cho việc dẫn xuất của các gradient đối với BDOF và PROF khi phương pháp thứ nhất (Fig.15) và phương pháp thứ hai (Fig.16) được áp dụng, một cách tương ứng. Trên Fig.15 và Fig.16, các hình tròn trắng biểu diễn các mẫu tham chiếu tại các vị trí nguyên, các hình tam giác biểu diễn các mẫu dự đoán phân số của khối hiện tại và các hình tròn theo mẫu hình biểu diễn các mẫu tham chiếu nguyên mà được sử dụng để lấp đầy vùng được mở rộng của khối hiện tại cho việc dẫn xuất gradient.

Theo cách bổ sung, theo các thiết kế BDOF và PROF sẵn có, thì việc đệm mẫu dự đoán được tiến hành tại các mức tạo mã khác nhau. Ví dụ, đối với BDOF, việc đệm được áp dụng dọc theo các ranh giới của mỗi khối con $sbWidth \times sbHeight$ mà tại đó $sbWidth = \min(CUWidth, 16)$ và $sbHeight = \min(CUHeight, 16)$. $CUWidth$ và $CUHeight$ là chiều rộng và chiều cao của một CU. Mặt khác, việc đệm của PROF luôn được áp dụng trên mức khối con 4×4 . Trong phần thảo luận ở trên, thì chỉ phương pháp đệm là thống nhất giữa BDOF và PROF trong khi các kích cỡ đệm khối con là vẫn khác nhau. Điều này cũng không thân thiện đối với cách triển khai phần cứng căn cứ vào việc các môđun khác nhau cần được triển khai cho các quy trình đệm của BDOF và PROF. Để đạt được một thiết kế thống nhất hơn, thì sáng chế đề xuất thống nhất kích cỡ đệm khối con của BDOF và PROF. Trong một hoặc nhiều phương án, sáng chế đề xuất áp dụng việc đệm mẫu dự đoán của BDOF tại mức 4×4 . Ví dụ, theo phương pháp này, CU được chia trước tiên thành nhiều khối con 4×4 ; sau việc bù chuyển động của mỗi khối con 4×4 , thì các mẫu được mở rộng dọc theo các ranh giới trên cùng/dưới cùng và bên trái/bên phải được đệm bằng cách sao chép các vị trí mẫu nguyên tương ứng. Fig.19A, Fig.19B, Fig.19C, và Fig.19D minh họa một ví dụ mà tại đó phương pháp đệm được đề xuất mà được áp dụng cho một CU BDOF 16×16 , mà tại đó các đường gạch biểu diễn các ranh giới khối con 4×4 và các dải màu xanh biểu diễn các mẫu được đệm của mỗi khối con 4×4 .

Fig.19A thể hiện phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16×16 , mà tại đó các đường gạch biểu diễn ranh giới khối con 4×4 trên cùng bên trái, theo sáng chế.

Fig.19B thể hiện phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16×16 , mà tại đó các đường gạch biểu diễn ranh giới khối con 4×4 trên cùng bên phải, theo sáng chế.

Fig.19C thể hiện phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, mà tại đó các đường gạch biểu diễn ranh giới khối con 4X4 dưới cùng bên trái 1960, theo sáng chế.

Fig.19D thể hiện phương pháp đệm được đề xuất mà được áp dụng cho CU BDOF 16X16, mà tại đó các đường gạch biểu diễn ranh giới khối con 4X4 dưới cùng bên phải 1980, theo sáng chế.

Cú pháp phát tín hiệu mức cao để cho phép/không cho phép BDOF, PROF và DMVR

Trong các thiết kế BDOF và PROF sẵn có, hai cờ khác nhau được phát tín hiệu trong tập hợp thông số chuỗi (SPS) để điều khiển việc cho phép/không cho phép của hai công cụ tạo mã theo kiểu tách biệt. Tuy nhiên, do tính tương tự giữa BDOF và PROF, nên rất mong muốn cho phép và/hoặc không cho phép BDOF và PROF từ mức cao bởi cùng một cờ điều khiển. Dựa trên việc xem xét như vậy, thì một cờ mới, vốn được gọi là `sps_bdof_prof_enabled_flag`, được giới thiệu tại SPS, như được thể hiện trong bảng 1. Như được thể hiện trong bảng 1, việc cho phép và không cho phép của BDOF chỉ phụ thuộc vào `sps_bdof_prof_enabled_flag`. Khi cờ bằng 1, thì BDOF được cho phép để tạo nội dung video trong chuỗi. Nếu không thì, khi `sps_bdof_prof_enabled_flag` bằng 0, thì BDOF sẽ không được áp dụng. Mặt khác, bên cạnh `sps_bdof_prof_enabled_flag`, thì cờ điều khiển afin mức SPS, tức là, `sps_affine_enabled_flag`, cũng được sử dụng để cho phép và không cho phép một cách có điều kiện PROF. Khi cả cờ `sps_bdof_prof_enabled_flag` và `sps_affine_enabled_flag` bằng 1, thì PROF được cho phép cho tất cả các khối tạo mã mà được tạo mã trong chế độ afin. Khi cờ `sps_bdof_prof_enabled_flag` bằng 1 và `sps_affine_enabled_flag` bằng 0, thì PROF không được cho phép.

Bảng 1 Bảng cú pháp SPS được sửa đổi với cờ cho phép/không cho phép BDOF/PROF được đề xuất

seq_parameter_set_rbsp() {	Descriptor
.....	
if(sps_temporal_mvp_enabled_flag)	
sps_sbtmvp_enabled_flag	u(1)
sps_amvr_enabled_flag	u(1)
sps_bdof_prof_enabled_flag	u(1)
sps_smvd_enabled_flag	u(1)
sps_affine_amvr_enabled_flag	u(1)
sps_dmvr_enabled_flag	u(1)
if(sps_bdof_prof_enabled_flag sps_dmvr_enabled_flag)	
sps_bdof_prof_dmvr_slice_present_flag	u(1)
sps_mmvd_enabled_flag	u(1)
sps_ism_enabled_flag	u(1)
sps_mrl_enabled_flag	u(1)
sps_mip_enabled_flag	u(1)
sps_cclm_enabled_flag	u(1)
.....	
}	

sps_bdof_prof_enabled_flag chỉ rõ liệu luồng quang học hai chiều và sự tinh chỉnh dự đoán với luồng quang học có được cho phép hay không. Khi **sps_bdof_prof_enabled_flag** bằng 0, thì cả luồng quang học hai chiều và sự tinh chỉnh dự đoán với luồng quang học đều không được cho phép. Khi **sps_bdof_prof_enabled_flag** bằng 1 và **sps_affine_enabled_flag** bằng 1, thì cả luồng quang học hai chiều và sự tinh chỉnh dự đoán với luồng quang học đều được cho phép. Nếu không thì (**sps_bdof_prof_enabled_flag** bằng 1 và **sps_affine_enabled_flag** bằng 0)

luồng quang học hai chiều được cho phép và sự tinh chỉnh dự đoán với luồng quang học không được cho phép.

sps_bdof_prof_dmvr_slice_preset_flag chỉ rõ khi cờ **slice_disable_bdof_prof_dmvr_flag** được phát tín hiệu tại mức lát. Khi cờ bằng 1, thì cú pháp **slice_disable_bdof_prof_dmvr_flag** được phát tín hiệu cho mỗi lát mà tham chiếu đến tập hợp thông số chuỗi hiện tại. Nếu không thì (khi **sps_bdof_prof_dmvr_slice_present_flag** bằng 0), cú pháp **slice_disabled_bdof_prof_dmvr_flag** sẽ không được phát tín hiệu tại mức lát. Khi cờ không được phát tín hiệu, thì nó được suy ra là 0.

Hơn nữa, khi cờ điều khiển BDOF và PROF mức SPS được đề xuất mà được sử dụng, thì cờ điều khiển tương ứng **no_bdof_constraint_flag** trong cú pháp thông tin ràng buộc chung cũng nên được sửa đổi bởi

general_constraint_info() {	Descriptor
...	
no_temporal_mvp_constraint_flag	u(1)
no_sbtmvp_constraint_flag	u(1)
no_amvr_constraint_flag	u(1)
no_bdof_prof_constraint_flag	u(1)
...	
while(!byte_aligned())	
gci_alignment_zero_bit	f(1)
}	

no_bdof_prof_constraint_flag bằng 1 chỉ rõ rằng **sps_bdof_prof_enabled_flag** sẽ bằng 0. **no_bdof_constraint_flag** bằng 0 không đặt ra ràng buộc.

Bên cạnh cú pháp BDOF/PROF SPS ở trên, thì sáng chế đề xuất giới thiệu một cờ điều khiển khác tại mức lát, tức là, `slice_disable_bdo_f_prof_dmvr_flag` được giới thiệu để không cho phép BDOF, PROF và DMVR. Cờ SPS `sps_bdo_f_prof_dmvr_slice_present_flag`, vốn được phát tín hiệu trong SPS khi hoặc một thành phần trong số trong số DMVR hoặc các cờ điều khiển mức BDOF SPS/PROF là đúng, được sử dụng để chỉ ra sự có mặt của `slice_disable_bdo_f_prof_dmvr_flag`. Nếu có mặt, thì `slice_disable_bdo_f_dmvr_flag` được phát tín hiệu. Bảng 2 minh họa bảng cú pháp phần đầu lát được sửa đổi sau khi cú pháp được đề xuất mà được áp dụng. Trong một phương án khác, sáng chế đề xuất vẫn sử dụng hai cờ điều khiển tại phần đầu lát để điều khiển theo kiểu tách biệt việc cho phép/không cho phép của BDOF và DMVR, và việc cho phép/không cho phép của PROF. Ví dụ, hai cờ được sử dụng trong phần đầu lát bởi phương pháp này: một cờ là `slice_disable_bdo_f_dmvr_slice_flag` được sử dụng để điều khiển bật/tắt của BDOF và DMVR và cờ khác `disable_prof_slice_flag` được sử dụng để điều khiển bật/tắt của chỉ PROF.

Bảng 2 Bảng cú pháp SPS được sửa đổi với cờ cho phép/không cho phép BDOF/PROF được đề xuất

<code>seq_parameter_set_rbsp() {</code>	
<code>if(sps_bdo_f_prof_dmvr_slice_present_flag)</code>	
<code>slice_disable_bdo_f_prof_dmvr_enabled_flag</code>	u(1)
<code>.....</code>	

Trong một phương án khác, sáng chế đề xuất điều khiển theo kiểu tách biệt BDOF và PROF bởi hai cờ SPS khác nhau. Ví dụ, hai cờ SPS tách biệt `sps_bdo_f_enable_flag` và `sps_prof_enable` được giới thiệu để cho phép/không cho phép hai công cụ theo kiểu tách biệt. Theo cách bổ sung, một cờ điều khiển mức cao `no_prof_constraint_flag`

cần được cộng vào bảng cú pháp `general_constrain_info()` không cho phép theo kiểu cường bức công cụ PROF.

<code>seq_parameter_set_rbsp() {</code>	Descriptor
<code>.....</code>	
<code>if(sps_temporal_mvp_enabled_flag)</code>	
<code> sps_sbtmvp_enabled_flag</code>	<code>u(1)</code>
<code> sps_amvr_enabled_flag</code>	<code>u(1)</code>
<code> sps_bdof_enabled_flag</code>	<code>u(1)</code>
<code> sps_prof_enabled_flag</code>	<code>u(1)</code>
<code> sps_smvd_enabled_flag</code>	<code>u(1)</code>
<code> sps_affine_amvr_enabled_flag</code>	<code>u(1)</code>
<code> sps_dmvr_enabled_flag</code>	<code>u(1)</code>
<code>if(sps_bdof_enabled_flag sps_dmvr_enabled_flag)</code>	
<code> sps_bdof_dmvr_slice_present_flag</code>	<code>u(1)</code>
<code> sps_mmvd_enabled_flag</code>	<code>u(1)</code>
<code> sps_isp_enabled_flag</code>	<code>u(1)</code>
<code> sps_mrl_enabled_flag</code>	<code>u(1)</code>
<code> sps_mip_enabled_flag</code>	<code>u(1)</code>
<code> sps_cclm_enabled_flag</code>	<code>u(1)</code>
<code>.....</code>	
<code>}</code>	

sps_bdof_enabled_flag chỉ rõ liệu luồng quang học hai chiều có được cho phép hay không. Khi `sps_bdof_` được cho phép_flag bằng 0, thì luồng quang học hai chiều không được cho phép. Khi `sps_bdof_enabled_flag` bằng 1, thì luồng quang học hai chiều được cho phép.

sps_prof_enabled_flag chỉ rõ liệu sự tinh chỉnh dự đoán với luồng quang học có được cho phép hay không. Khi `sps_prof_enabled_flag` bằng 0, thì sự tinh chỉnh dự đoán

với luồng quang học không được cho phép. Khi `sps_prof_enabled_flag` bằng 1, thì sự tinh chỉnh dự đoán với luồng quang học được cho phép.

<code>general_constraint_info() {</code>	Descriptor
...	
<code>no_temporal_mvp_constraint_flag</code>	<code>u(1)</code>
<code>no_sbtmvp_constraint_flag</code>	<code>u(1)</code>
<code>no_amvr_constraint_flag</code>	<code>u(1)</code>
<code>no_bdof_constraint_flag</code>	<code>u(1)</code>
<code>no_prof_constraint_flag</code>	<code>u(1)</code>
...	
<code>while(!byte_aligned())</code>	
<code>gci_alignment_zero_bit</code>	<code>f(1)</code>
<code>}</code>	

`no_prof_constraint_flag` bằng 1 chỉ rõ rằng `sps_prof_enabled_flag` sẽ bằng 0. `no_prof_constraint_flag` bằng 0 không đặt ra ràng buộc.

Tại mức lát, trong một hoặc nhiều phương án của sáng chế, sáng chế đề xuất giới thiệu một cờ điều khiển khác tại mức lát, tức là, `slice_disable_bdof_prof_dmvr_flag` được giới thiệu để không cho phép BDOF, PROF và DMVR cùng nhau. Trong một phương án khác, sáng chế đề xuất cộng hai cờ tách biệt, tức là, `slice_disable_bdof_dmvr_flag` và `slice_disable_prof_flag`, tại mức lát. Cờ thứ nhất (tức là, `slice_disable_bdof_dmvr_flag`) được sử dụng để bật/tắt theo kiểu thích ứng BDOF và DMVR cho một lát trong khi cờ thứ hai (tức là, `slice_disable_prof_flag`) được sử dụng để điều khiển việc cho phép và không cho phép của công cụ PROF tại mức lát. Theo cách bổ sung, khi phương pháp thứ hai được áp dụng, thì cờ `slice_disable_bdof_dmvr_flag` chỉ cần được phát tín hiệu khi hoặc cờ BDOF SPS hoặc

DMVR SPS được cho phép và cờ chỉ cần được phát tín hiệu khi cờ SPS PROF được cho phép.

Tại phiên họp JVET lần thứ 16, phần đầu ảnh đã được chấp thuận vào trong bản thảo VVC. Phần đầu ảnh được phát tín hiệu một lần trên mỗi lát mà các phần tử cú pháp nằm trước đó trong phần đầu lát và không thay đổi từ lát thành lát.

Dựa trên phần đầu ảnh được chấp thuận, trong một hoặc nhiều phương án của sáng chế, sáng chế đề xuất điều khiển các cờ điều khiển BDOF, DMVR và PROF từ phần đầu lát thành phần đầu ảnh. Ví dụ, trong phương pháp được đề xuất, thì ba cờ điều khiển khác nhau `sps_dmvr_picture_header_present_flag`, `sps_bdof_picture_header_present_flag` và `sps_prof_picture_header_present_flag` được phát tín hiệu trong SPS. Khi một cờ trong số ba cờ được phát tín hiệu dưới dạng đúng, thì một cờ điều khiển bổ sung sẽ được phát tín hiệu trong phần đầu ảnh để chỉ ra công cụ tương ứng (tức là, DMVR, BDOF và PROF) đánh được cho phép hoặc không được cho phép cho các lát tham chiếu đến phần đầu ảnh. Các phần tử cú pháp được đề xuất sẽ được chỉ rõ như sau.

<code>seq_parameter_set_rbsp() {</code>	Descriptor
<code>.....</code>	
<code>if(sps_temporal_mvp_enabled_flag)</code>	
<code> sps_sbtmvp_enabled_flag</code>	<code>u(1)</code>
<code> sps_amvr_enabled_flag</code>	<code>u(1)</code>
<code> sps_bdof_enabled_flag</code>	<code>u(1)</code>
<code> sps_prof_enabled_flag</code>	<code>u(1)</code>
<code> sps_smvd_enabled_flag</code>	<code>u(1)</code>
<code> sps_affine_amvr_enabled_flag</code>	<code>u(1)</code>
<code> sps_dmvr_enabled_flag</code>	<code>u(1)</code>
<code>if(sps_dmvr_enabled_flag)</code>	

sps_dmvr_picture_header_present_flag	u(1)
if(sps_bdof_enabled_flag)	
sps_bdof_picture_header_present_flag	u(1)
sps_mmvd_enabled_flag	u(1)
sps_isp_enabled_flag	u(1)
sps_mrl_enabled_flag	u(1)
sps_mip_enabled_flag	u(1)
sps_cclm_enabled_flag	u(1)
.....	
sps_affine_enabled_flag	u(1)
if(sps_affine_enabled_flag) {	
sps_affine_type_flag	u(1)
sps_affine_amvr_enabled_flag	u(1)
sps_affine_prof_enabled_flag	u(1)
if(sps_affine_prof_enabled_flag)	
sps_prof_picture_header_present_flag	u(1)
}	
}	

sps_dmvr_picture_header_preset_flag chỉ rõ liệu cờ picture_disable_dmvr_flag có được phát tín hiệu tại phần đầu ảnh hay không. Khi cờ bằng 1, thì cú pháp picture_disable_dmvr_flag được phát tín hiệu cho mỗi ảnh mà tham chiếu đến tập hợp thông số chuỗi hiện tại. Nếu không thì, cú pháp picture_disable_dmvr_flag sẽ không được phát tín hiệu tại phần đầu ảnh. Khi cờ không được phát tín hiệu, thì nó được suy ra là 0.

sps_bdof_picture_header_preset_flag chỉ rõ liệu cờ picture_disable_bdof_flag có được phát tín hiệu tại phần đầu ảnh hay không. Khi cờ bằng 1, thì cú pháp picture_disable_bdof_flag được phát tín hiệu cho mỗi ảnh mà tham chiếu đến tập hợp thông số chuỗi hiện tại. Nếu không thì, cú pháp picture_disable_bdof_flag sẽ không

được phát tín hiệu tại phần đầu ảnh. Khi cờ không được phát tín hiệu, thì nó được suy ra là 0.

sps_prof_picture_header_preset_flag chỉ rõ liệu cờ **picture_disable_prof_flag** được phát tín hiệu tại phần đầu ảnh. Khi cờ bằng 1, thì cú pháp **picture_disable_prof_flag** được phát tín hiệu cho mỗi ảnh mà tham chiếu đến tập hợp thông số chuỗi hiện tại. Nếu không thì, cú pháp **picture_disable_prof_flag** sẽ không được phát tín hiệu tại phần đầu ảnh. Khi cờ không được phát tín hiệu, thì nó được suy ra là 0.

picture_header_rbps() {	Descriptor
if(sps_dmvr_picture_header_present_flag)	
picture_disable_dmvr_flag	u(1)
if(sps_bdof_picture_header_present_flag)	
picture_disable_bdof_flag	u(1)
if(sps_prof_picture_header_present_flag)	
picture_disable_prof_flag	u(1)
}	

picture_disable_dmvr_flag chỉ rõ liệu công cụ dmvr có được cho phép cho các lát mà tham chiếu đến phần đầu ảnh hiện tại hay không. Khi cờ bằng 1, thì công cụ dmvr được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại. Nếu không thì, công cụ dmvr không được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại. Khi cờ không có mặt, thì cờ được suy ra là 0.

picture_disable_bdof_flag chỉ rõ liệu công cụ bdof có được cho phép cho các lát mà tham chiếu đến phần đầu ảnh hiện tại hay không. Khi cờ bằng 1, thì công cụ bdof được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại. Nếu không thì, công cụ bdof không được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại.

picture_disable_prof_flag chỉ rõ liệu công cụ prof có được cho phép cho các lát mà tham chiếu đến phần đầu ảnh hiện tại hay không. Khi cờ bằng 1, thì công cụ prof được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại. Nếu không thì, công cụ prof không được cho phép cho các lát tham chiếu đến phần đầu ảnh hiện tại.

Fig.12 thể hiện phương pháp của BDOF, PROF, và DMVR. Phương pháp có thể, ví dụ, được áp dụng cho bộ giải mã.

Trong bước 1210, bộ giải mã có thể nhận ba cờ điều khiển trong tập hợp thông số chuỗi (SPS). Cờ điều khiển thứ nhất chỉ ra liệu BDOF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ hai chỉ ra liệu PROF có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không. Cờ điều khiển thứ ba chỉ ra liệu DMVR có được cho phép để giải mã các khối video trong chuỗi video hiện tại hay không.

Trong bước 1212, bộ giải mã có thể nhận cờ có mặt thứ nhất trong SPS khi cờ điều khiển thứ nhất là đúng, cờ có mặt thứ hai trong SPS khi cờ điều khiển thứ hai là đúng và cờ có mặt thứ ba trong SPS khi cờ điều khiển thứ ba là đúng.

Trong bước 1214, bộ giải mã có thể nhận cờ điều khiển ảnh thứ nhất trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ nhất trong SPS chỉ ra rằng BDOF không được cho phép cho các khối video ở trong ảnh.

Trong bước 1216, bộ giải mã có thể nhận cờ điều khiển ảnh thứ hai trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ hai trong SPS chỉ ra rằng PROF không được cho phép cho các khối video ở trong ảnh.

Trong bước 1218, bộ giải mã có thể nhận cờ điều khiển ảnh thứ ba trong phần đầu ảnh của mỗi ảnh khi cờ có mặt thứ ba trong SPS chỉ ra rằng DMVR không được cho phép cho các khối video ở trong ảnh.

Việc kết thúc sớm của PROF dựa trên sự chênh lệch MV điểm điều khiển

Theo thiết kế PROF hiện tại, PROF luôn được gọi khi một khối tạo mã được dự đoán bởi chế độ afin. Tuy nhiên, như được chỉ ra trong phương trình (6) và (7), các MV khối con của một khối afin được dẫn xuất từ các MV điểm điều khiển. Do đó, khi sự chênh lệch giữa các MV điểm điều khiển là tương đối nhỏ, thì các MV tại mỗi vị trí mẫu là phải nhất quán. Trong trường hợp như vậy, thì lợi ích của áp dụng PROF có thể rất bị hạn chế. Do đó, để giảm thêm nữa độ phức tạp tính toán trung bình của PROF, thì sáng chế đề xuất bỏ qua theo kiểu thích ứng sự tinh chỉnh mẫu dựa trên PROF dựa trên sự chênh lệch MV tối đa giữa MV đảo mẫu và MV đảo khối con ở trong một khối con 4x4. Bởi vì các giá trị của sự chênh lệch MV PROF của các mẫu bên trong một khối con 4x4 là đối xứng quanh trung tâm khối con, nên sự chênh lệch MV PROF ngang và dọc tối đa có thể được tính toán dựa trên phương trình (10) dưới dạng

$$\begin{aligned}\Delta v_x^{max} &= 6 * (c + d) \\ \Delta v_y^{max} &= 6 * (e + f)\end{aligned}\tag{19}$$

Theo sáng chế, các chuẩn đo khác nhau có thể được sử dụng trong việc xác định nếu sự chênh lệch MV đủ nhỏ để bỏ qua quy trình PROF.

Trong một ví dụ, dựa trên phương trình (19), quy trình PROF có thể được bỏ qua khi tổng của sự chênh lệch MV ngang tối đa tuyệt đối và sự chênh lệch MV dọc tối đa tuyệt đối nhỏ hơn một được định rõ trước ngưỡng, tức là,

$$|\Delta v_x^{max}| + |\Delta v_y^{max}| \leq thres\tag{20}$$

Trong một ví dụ khác, nếu giá trị tối đa của $|\Delta v_x^{max}|$ và $|\Delta v_y^{max}|$ không lớn hơn ngưỡng, thì quy trình PROF có thể được bỏ qua.

$$\text{MAX}(|\Delta v_x^{max}|, |\Delta v_y^{max}|) \leq \text{thres} \quad (21)$$

Mà tại đó $\text{MAX}(a, b)$ là hàm trả lại giá trị lớn hơn giữa các giá trị đầu vào a và b .

Bên cạnh hai ví dụ ở trên, thì mục đích của sáng chế cũng có thể áp dụng được cho các trường hợp khi các chuẩn đo khác được sử dụng trong việc xác định nếu sự chênh lệch MV đủ nhỏ để bỏ qua quy trình PROF.

Trong phương pháp ở trên, thì PROF được bỏ qua dựa trên độ lớn của sự chênh lệch MV. Mặt khác, bên cạnh sự chênh lệch MV, thì sự tinh chỉnh mẫu PROF cũng được tính toán dựa trên thông tin gradient cục bộ tại mỗi sự định vị mẫu trong một khối chuyển động được bù. Đối với các khối dự đoán mà chứa ít chi tiết tần số cao hơn (ví dụ, khu vực phẳng), thì các giá trị gradient có xu hướng là nhỏ sao cho các giá trị của các sự tinh chỉnh mẫu được dẫn xuất là nhỏ. Xem xét đến điều này, thì theo một phương án khác, sáng chế đề xuất chỉ áp dụng PROF cho các mẫu dự đoán của các khối mà chứa đủ thông tin tần số cao.

Các chuẩn đo khác nhau có thể được sử dụng trong việc xác định nếu khối chứa đủ thông tin tần số cao để quy trình PROF đáng để được gọi cho khối. Trong một ví dụ, việc quyết định được thực hiện dựa trên độ lớn trung bình (tức là giá trị tuyệt đối) của các gradient của các mẫu ở trong khối dự đoán. Nếu độ lớn trung bình nhỏ hơn một ngưỡng, thì sau đó khối dự đoán được phân loại là khu vực phẳng và PROF sẽ không được áp dụng; nếu không thì, khối dự đoán được xem xét để chứa đủ các chi tiết tần số cao mà tại đó PROF vẫn có thể áp dụng được. Trong một ví dụ khác, độ lớn tối đa của các gradient của các mẫu ở trong khối dự đoán có thể được sử dụng. Nếu độ lớn tối đa

nhỏ hơn một ngưỡng, thì PROF cần được bỏ qua cho khối. Trong một ví dụ khác nữa, sự chênh lệch giữa giá trị mẫu tối đa và giá trị mẫu tối thiểu, $I_{max} - I_{min}$, của khối dự đoán có thể được sử dụng để xác định nếu PROF cần được áp dụng cho khối. Nếu giá trị chênh lệch nhỏ hơn ngưỡng, thì PROF cần được bỏ qua cho khối. Cần lưu ý rằng mục đích của sáng chế cũng có thể áp dụng được cho các trường hợp mà tại đó một số chuẩn đo khác được sử dụng trong việc xác định nếu khối đã cho chứa đủ thông tin tần số cao hay không.

Xử lý sự tương tác giữa PROF và LIC cho chế độ afin

Bởi vì các mẫu được tái tạo lân cận (tức là, khuôn mẫu) của khối hiện tại được sử dụng bởi LIC để dẫn xuất các thông số mô hình tuyến tính, nên việc giải mã của một LIC khối tạo mã phụ thuộc vào việc tái tạo đầy đủ của các mẫu lân cận của nó. Do sự phụ thuộc lẫn nhau như vậy, nên đối với các cách triển khai phần cứng thực tế, thì LIC cần được thực hiện trong giai đoạn tái tạo mà tại đó các mẫu được tái tạo lân cận trở nên khả dụng cho việc dẫn xuất thông số LIC. Bởi vì việc tái tạo khối phải được thực hiện một cách tuần tự (tức là, từng khối một), nên thông lượng (tức là, lượng công việc mà có thể được thực hiện song song trên mỗi đơn vị thời gian) là một vấn đề quan trọng để xem xét khi áp dụng chung các phương pháp tạo mã khác cho các khối tạo mã LIC. Trong phần này, hai phương pháp được đề xuất để xử lý sự tương tác khi cả PROF và LIC được cho phép cho chế độ afin.

Trong phương án thứ nhất, sáng chế đề xuất áp dụng theo cách dành riêng chế độ PROF và chế độ LIC cho một khối tạo mã afin. Như được thảo luận trước đây, trong thiết kế sẵn có, PROF được áp dụng một cách rõ ràng cho tất cả afin các khối mà không phát tín hiệu trong khi một cờ LIC được phát tín hiệu hoặc được thừa hưởng tại mức khối tạo mã để chỉ ra liệu chế độ LIC có được áp dụng cho một khối afin hay không.

Theo phương pháp sáng chế đề xuất áp dụng một cách có điều kiện PROF dựa trên giá trị của cờ LIC của một khối afin. Khi cờ bằng 1, thì chỉ LIC được áp dụng bằng cách điều chỉnh các mẫu dự đoán của toàn bộ khối tạo mã dựa trên trọng số và dịch vị LIC. Nếu không thì (tức là, cờ LIC bằng 0), PROF được áp dụng cho khối tạo mã afin để tinh chỉnh các mẫu dự đoán của mỗi khối con dựa trên mô hình luồng quang học.

Fig.18A minh họa một lưu đồ làm ví dụ của quy trình giải mã dựa trên phương pháp được đề xuất mà tại đó PROF và LIC không được chấp nhận để được áp dụng đồng thời.

Fig.18A thể hiện sự minh họa của quy trình giải mã dựa trên phương pháp được đề xuất mà tại đó PROF và LIC không được chấp nhận, theo sáng chế. Quy trình giải mã 1820 gồm bước xác định 1822, LIC 1824, và PROF 1826. Bước xác định 1822 xác định liệu cờ LIC có bật hay không và bước tiếp theo được thực hiện theo việc xác định đó. LIC 1824 là ứng dụng của LIC nếu cờ LIC được thiết lập. PROF 1826 là ứng dụng của PROF nếu cờ LIC không được thiết lập.

Trong phương án thứ hai, sáng chế đề xuất áp dụng LIC sau PROF để tạo ra các mẫu dự đoán của một khối afin. Ví dụ, sau khi việc bù chuyển động afin dựa trên khối con được thực hiện, thì các mẫu dự đoán được tinh chỉnh dựa trên sự tinh chỉnh mẫu PROF; sau đó, LIC được tiến hành bằng cách áp dụng một cặp trọng số và dịch vị (như được dẫn xuất từ khuôn mẫu và các mẫu tham chiếu của nó) cho các mẫu dự đoán PROF-được điều chỉnh để thu nhận các mẫu dự đoán cuối cùng của khối, như được minh họa dưới dạng

$$P[\mathbf{x}] = \alpha * (P_r[\mathbf{x} + \mathbf{v}] + \Delta I[\mathbf{x}]) + \beta \quad (22)$$

mà tại đó $P_r[\mathbf{x} + \mathbf{v}]$ là khối tham chiếu của khối hiện tại được chỉ ra bởi véc-tơ chuyển động $\mathbf{v}$; α và β là trọng số và dịch vị LIC; $P[\mathbf{x}]$ là khối dự đoán cuối cùng; $\Delta[\mathbf{x}]$ là sự tinh chỉnh PROF như được dẫn xuất trong (17).

Fig.18B thể hiện sự minh họa của quy trình giải mã mà tại đó PROF và LIC được áp dụng, theo sáng chế. Quy trình giải mã 1860 gồm việc bù chuyển động afin 1862, việc dẫn xuất thông số LIC 1864, PROF 1866, và việc điều chỉnh mẫu LIC 1868. Việc bù chuyển động afin 1862 áp dụng chuyển động afin và là đầu vào cho việc dẫn xuất thông số LIC 1864 và PROF 1866. Việc dẫn xuất thông số LIC 1864 được áp dụng để dẫn xuất các thông số LIC. PROF 1866 là PROF đang được áp dụng. Việc điều chỉnh mẫu LIC 1868 là các thông số trọng số và dịch vị LIC đang được tổ hợp với PROF.

Fig.18B minh họa tiến trình công việc giải mã làm ví dụ khi phương pháp thứ hai được áp dụng. Như được thể hiện trên Fig.18B, bởi vì LIC sử dụng khuôn mẫu (tức là, các mẫu được tái tạo lân cận) để tính toán mô hình tuyến tính LIC, nên các thông số LIC có thể được dẫn xuất ngay lập tức ngay khi các mẫu được tái tạo lân cận trở nên khả dụng. Điều này có nghĩa là sự tinh chỉnh PROF và việc dẫn xuất thông số LIC có thể được tiến hành tại cùng thời điểm.

Khối lượng và dịch vị LIC (tức là, α và β) và sự tinh chỉnh PROF (tức là, $\Delta[\mathbf{x}]$) là các số động chung. Đối với các cách triển khai phần cứng thân thiện, thì các phép toán số động đó thường được triển khai dưới dạng phép nhân với một giá trị nguyên được theo sau bởi phép toán dịch chuyển sang phải theo số lượng các bit. Trong thiết kế LIC và PROF sẵn có, vì hai công cụ được thiết kế theo kiểu tách biệt, nên hai dịch chuyển sang phải khác nhau, theo N_{LIC} bit và N_{PROF} bit một cách tương ứng, được áp dụng tại hai giai đoạn.

Theo phương án thứ ba, để cải thiện độ lợi tạo mã trong trường hợp PROF và LIC được áp dụng chung cho khối tạo mã afin, thì sáng chế đề xuất áp dụng các sự điều chỉnh mẫu dựa trên LIC và dựa trên PROF tại độ chính xác cao. Việc này được thực hiện bằng cách tổ hợp hai phép toán dịch chuyển sang phải của chúng thành một và áp dụng nó tại lúc kết thúc để dẫn xuất các mẫu dự đoán cuối cùng (như được thể hiện trong (12)) của khối hiện tại.

Giải quyết vấn đề tràn phép nhân khi tổ hợp PROF với việc dự đoán có trọng số và việc dự đoán đôi với trọng số mức CU (Bi-prediction with CU-level Weight, BCW)

Theo thiết kế PROF hiện tại trong bản thảo làm việc VVC, thì PROF có thể được áp dụng chung với việc dự đoán có trọng số (Weighted Prediction, WP). Ví dụ, khi chúng được tổ hợp, thì tín hiệu dự đoán của một CU afin sẽ được tạo ra bởi các thủ tục sau đây:

Thứ nhất, đối với mỗi mẫu tại vị trí (x,y) , thì tính toán sự tinh chỉnh dự đoán L0 $\Delta I_0(x,y)$ dựa trên PROF và cộng sự tinh chỉnh vào mẫu dự đoán L0 gốc $I_0(x,y)$, tức là,

$$\begin{aligned} \Delta I_0(x,y) &= (g_{h0}(x,y) \cdot \Delta v_{x0}(x,y) + g_{v0}(x,y) \cdot \Delta v_{y0}(x,y) + 1) \gg 1 \\ I'_0(x,y) &= I_0(x,y) + \Delta I_0(x,y) \end{aligned} \quad (23)$$

mà tại đó $I'_0(x,y)$ là mẫu được tinh chỉnh; $g_{h0}(x,y)$ và $g_{v0}(x,y)$ và $\Delta v_{x0}(x,y)$ và $\Delta v_{y0}(x,y)$ là các gradient ngang/dọc L0 và các sự tinh chỉnh chuyển động ngang/dọc L0 tại vị trí (x,y) .

Thứ hai, đối với mỗi mẫu tại vị trí (x,y) , thì tính toán sự tinh chỉnh dự đoán L1 $\Delta I_1(x,y)$ dựa trên PROF và cộng sự tinh chỉnh vào các mẫu dự đoán L1 gốc $I_1(x,y)$, tức là,

$$\begin{aligned} \Delta I_1(x,y) &= (g_{h1}(x,y) \cdot \Delta v_{x1}(x,y) + g_{v1}(x,y) \cdot \Delta v_{y1}(x,y) + 1) \gg 1 \\ I'_1(x,y) &= I_1(x,y) + \Delta I_1(x,y) \end{aligned} \quad (24)$$

mà tại đó $I'_1(x,y)$ là mẫu được tinh chỉnh; $g_{h1}(x,y)$ và $g_{v1}(x,y)$ và $\Delta v_{x1}(x,y)$ và $\Delta v_{y1}(x,y)$ là các gradient ngang/dọc L1 và các sự tinh chỉnh chuyển động ngang/dọc L1 tại vị trí (x,y) .

Thứ ba, tổ hợp các mẫu dự đoán L0 và L1 được tinh chỉnh, tức là,

$$I_{bi}(x,y) = (W_0 \cdot I'_0(x,y) + W_1 \cdot I'_1(x,y) + Offset) \gg shift \quad (25)$$

mà tại đó W_0 và W_1 là trọng số WP và BCW; $shift$ và $Offset$ là dịch vị và dịch chuyển sang phải mà được áp dụng cho trung bình có trọng số của các tín hiệu dự đoán L0 và L1 cho việc dự đoán đôi cho WP và BCW. Ở đây, các thông số cho WP gồm W_0 và $Offset$ trong khi các thông số cho BCW gồm W_0 và W_1 và $shift$.

Như có thể thấy được từ các phương trình ở trên, do sự tinh chỉnh đảo mẫu, tức là, $\Delta I_0(x,y)$ và $\Delta I_1(x,y)$, nên các mẫu dự đoán sau PROF (tức là, $I'_0(x,y)$ và $I'_1(x,y)$) sẽ có một phạm vi linh động tăng lên so với các mẫu dự đoán gốc (tức là, $I_0(x,y)$ và $I_1(x,y)$). Căn cứ vào việc các mẫu dự đoán được tinh chỉnh sẽ được nhân với các hệ số trọng số WP và BCW, thì điều này sẽ tăng độ dài của số nhân được yêu cầu. Ví dụ, dựa trên thiết kế hiện tại, khi độ sâu bit tạo mã bên trong có phạm vi từ 8 đến 12-bit, thì phạm vi linh động của tín hiệu dự đoán $I_0(x,y)$ và $I_1(x,y)$ là 16-bit. Nhưng, sau PROF, thì phạm vi linh động của tín hiệu dự đoán $I'_0(x,y)$ và $I'_1(x,y)$ là 17-bit. Do đó, khi PROF được áp dụng, thì có thể có tiềm năng là gây ra vấn đề tràn phép nhân 16-bit. Để xử lý vấn đề tràn này, thì nhiều phương pháp được đề xuất trong phần sau đây:

Thứ nhất, trong phương pháp thứ nhất, sáng chế đề xuất không cho phép WP và BCW khi PROF được áp dụng cho một CU afin.

Thứ hai, trong phương pháp thứ hai, sáng chế đề xuất áp dụng một phép toán cắt xén cho các sự tinh chỉnh mẫu được dẫn xuất trước khi cộng vào các mẫu dự đoán gốc

sao cho phạm vi linh động của các mẫu dự đoán được tinh chỉnh $I'_0(x,y)$ và $I'_1(x,y)$ sẽ có cùng độ sâu bit linh động như độ sâu bit của các mẫu dự đoán gốc $I_0(x,y)$ và $I_1(x,y)$. Ví dụ, nhờ phương pháp này, thì sự tinh chỉnh mẫu $\Delta I_0(x,y)$ và $\Delta I_1(x,y)$ sẽ trong (23) và (24) sẽ được sửa đổi bằng cách giới thiệu một phép toán cắt xén như được miêu tả dưới dạng:

$$\Delta I_0(x,y) = \text{clip3}(-2^{dI-1}, 2^{dI-1} - 1, \Delta I_0(x,y))$$

$$\Delta I_1(x,y) = \text{clip3}(-2^{dI-1}, 2^{dI-1} - 1, \Delta I_1(x,y))$$

mà tại đó $dI = dI_{base} + \max(0, BD - 12)$, mà tại đó BD là độ sâu bit tạo mã bên trong; dI_{base} là độ sâu bit cơ sở. Trong một hoặc nhiều phương án, sáng chế đề xuất thiết lập giá trị của dI_{base} là 14. Trong một phương án khác, sáng chế đề xuất thiết lập giá trị là 13. Trong một phương án khác, sáng chế đề xuất thiết lập một cách trực tiếp giá trị của dI là cố định. Trong một ví dụ, sáng chế đề xuất thiết lập giá trị của dI là 13, tức là, sự tinh chỉnh mẫu sẽ được cắt xén thành phạm vi $[-4096, 4095]$. Trong một ví dụ khác, sáng chế đề xuất thiết lập giá trị của dI là 14, tức là, sự tinh chỉnh mẫu sẽ được cắt xén thành phạm vi $[-8192, 8191]$.

Thứ nhất, trong phương pháp thứ ba, sáng chế đề xuất cắt xén một cách trực tiếp các mẫu dự đoán được tinh chỉnh thay vì cắt xén các sự tinh chỉnh mẫu sao cho các mẫu được tinh chỉnh cùng phạm vi linh động như phạm vi của các mẫu dự đoán gốc. Ví dụ, nhờ phương pháp thứ ba, các mẫu L0 và L1 được tinh chỉnh sẽ là

$$I'_0(x,y) = \text{clip3}(-2^{dR}, 2^{dR} - 1, I_0(x,y) + \Delta I_0(x,y))$$

$$I'_1(x,y) = \text{clip3}(-2^{dR}, 2^{dR} - 1, I_1(x,y) + \Delta I_1(x,y))$$

mà tại đó $dR = 16 + \max(0, BD-12)$ (hoặc theo cách tương đương là $\max(16, BD+4)$), mà tại đó BD là độ sâu bit tạo mã bên trong.

Thứ hai, trong phương pháp thứ tư, sáng chế đề xuất áp dụng các dịch chuyển sang phải nhất định cho các mẫu dự đoán L0 và L1 được tinh chỉnh trước WP và BCW; sau đó các mẫu dự đoán cuối cùng được điều chỉnh đến các độ chính xác gốc bằng cách bổ sung các dịch chuyển sang trái. Ví dụ, các mẫu dự đoán cuối cùng được dẫn xuất dưới dạng

$$I_{bi}(x,y) = (W_0 \cdot (I'_0(x,y) \gg nb) + W_1 \cdot (I'_1(x,y) \gg nb) + Offset) \gg (shift - nb)$$

mà tại đó nb là số lượng các dịch chuyển bit bổ sung mà được áp dụng vốn có thể được xác định dựa trên phạm vi linh động tương ứng của các sự tinh chỉnh mẫu PROF.

Thứ ba, trong phương pháp thứ năm, sáng chế đề xuất chia mỗi phép nhân của mẫu dự đoán L0/L1 với trọng số WP/BCW tương ứng trong (25) thành hai phép nhân và cả hai phép nhân không vượt quá 16-bit, như được mô tả dưới dạng

$$I_{bi}(x,y) = (W_0 \cdot I_0(x,y) + W_0 \cdot \Delta I_0(x,y) + W_1 \cdot I_1(x,y) + W_1 \cdot \Delta I_1(x,y) + Offset) \gg shift$$

Fig.20 thể hiện môi trường điện toán 2010 được ghép nối với giao diện người dùng 2060. Môi trường điện toán 2010 có thể là một phần của máy chủ xử lý dữ liệu. Môi trường điện toán 2010 gồm bộ xử lý 2020, bộ nhớ 2040, và giao diện I/O 2050.

Bộ xử lý 2020 thường điều khiển các hoạt động tổng thể của môi trường điện toán 2010, như là các hoạt động được kết hợp với sự hiển thị, thu thập dữ liệu, truyền thông dữ liệu, và xử lý hình ảnh. Bộ xử lý 2020 có thể gồm một hoặc nhiều bộ xử lý để thực thi các lệnh để thực hiện tất cả hoặc một số bước trong số các bước trong các phương pháp được mô tả ở trên. Hơn thế nữa, bộ xử lý 2020 có thể gồm một hoặc nhiều môđun mà tạo thuận tiện cho sự tương tác giữa bộ xử lý 2020 và các thành phần khác. Bộ xử lý có thể là đơn vị xử lý trung tâm (Central Processing Unit, CPU), bộ vi xử lý, máy chip đơn lẻ, GPU, hoặc tương tự.

Bộ nhớ 2040 được tạo cấu hình để lưu trữ các loại dữ liệu khác nhau để hỗ trợ các hoạt động của môi trường điện toán 2010. Bộ nhớ 2040 có thể gồm phần mềm được xác định trước 2042. Các ví dụ về dữ liệu như vậy bao gồm các lệnh cho các ứng dụng hoặc các phương pháp bất kỳ mà được vận hành trên môi trường điện toán 2010, các bộ dữ liệu video, dữ liệu hình ảnh, v.v. Bộ nhớ 2040 có thể được triển khai bằng cách sử dụng loại bất kỳ trong số các thiết bị bộ nhớ khả biến hoặc bất khả biến, hoặc tổ hợp của chúng, như là bộ nhớ truy cập ngẫu nhiên tĩnh (Static Random Access Memory, SRAM), bộ nhớ chỉ đọc lập trình được, xóa được bằng điện (Electrically Erasable Programmable Read-Only Memory, EEPROM), bộ nhớ chỉ đọc lập trình được, xóa được (Erasable Programmable Read-Only Memory, EPROM), bộ nhớ chỉ đọc lập trình được (Programmable Read-Only Memory, PROM), bộ nhớ chỉ đọc (Read-Only Memory, ROM), bộ nhớ từ, bộ nhớ tác động nhanh, đĩa từ hoặc quang.

Giao diện I/O 2050 cung cấp giao diện giữa bộ xử lý 2020 và các môđun giao diện ngoại vi, như là bàn phím, bánh xe nhấp chuột, các nút, và tương tự. Các nút có thể gồm nhưng không bị giới hạn ở, nút Home, nút khởi động quét, và nút dừng quét. Giao diện I/O 2050 có thể được ghép nối với bộ mã hóa và bộ giải mã.

Trong một số phương án, sáng chế cũng đề xuất phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp bao gồm nhiều chương trình, như là được bao gồm trong bộ nhớ 2040, có thể thực thi được bởi bộ xử lý 2020 trong môi trường điện toán 2010, để thực hiện các phương pháp được mô tả ở trên. Ví dụ, phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp có thể là ROM, RAM, CD-ROM, băng từ, đĩa mềm, thiết bị lưu trữ dữ liệu quang hoặc tương tự.

Phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp đã lưu trữ nhiều chương trình để thực thi bởi thiết bị điện toán có một hoặc nhiều bộ xử lý, mà tại đó

nhiều chương trình, khi được thực thi bởi một hoặc nhiều bộ xử lý, sẽ làm cho thiết bị điện toán thực hiện các phương pháp được mô tả ở trên để dự đoán chuyển động.

Trong một số phương án, môi trường điện toán 2010 có thể được triển khai với một hoặc nhiều mạch tích hợp chuyên dụng (Application Specific Integrated Circuit, ASIC), bộ xử lý tín hiệu (Digital Signal Processor, DSP), thiết bị xử lý tín hiệu (Digital Signal Processing Device, DSPD), thiết bị logic lập trình được (Programmable Logic Device, PLD), mảng cổng lập trình được dạng trường (Field Programmable Gate Array, FPGA), bộ điều khiển, bộ vi điều khiển, bộ vi xử lý, hoặc các thành phần điện tử khác, để thực hiện các phương pháp được mô tả trên đây.

Phần mô tả của sáng chế được trình bày nhằm mục đích minh họa, và không được dự định để trình bày đầy đủ hoặc giới hạn sáng chế. Nhiều sự sửa đổi, biến đổi, và cách triển khai thay thế sẽ trở nên rõ ràng với người có hiểu biết trung bình trong lĩnh vực kỹ thuật tương ứng có lợi ích trong việc giảng dạy được trình bày trong các phần mô tả được nêu ở trên và các hình vẽ liên kết.

Các ví dụ đã được chọn và được mô tả để giải thích rõ các nguyên tắc của sáng chế và để cho phép người có hiểu biết trung bình trong lĩnh vực kỹ thuật tương ứng hiểu sáng chế ở các cách triển khai khác nhau và để sử dụng tốt nhất các nguyên tắc cơ bản và các cách triển khai khác nhau với các biến thể khác nhau như được làm phù hợp với mục đích sử dụng cụ thể đã được dự tính. Do đó, cần hiểu rằng phạm vi của sáng chế không bị giới hạn ở các ví dụ cụ thể của các cách triển khai được bộc lộ và rằng các biến thể và các cách triển khai khác được dự tính để được bao gồm trong phạm vi của phần bộc lộ hiện tại.

YÊU CẦU BẢO HỘ

1. Phương pháp tinh chỉnh dự đoán với luồng quang học (Prediction Refinement with Optical Flow, PROF) để giải mã tín hiệu video, bao gồm các bước:

thu nhận, tại bộ giải mã, ảnh tham chiếu I được kết hợp với khối video;

thu nhận, tại bộ giải mã, các mẫu dự đoán ban đầu của khối video từ khối tham chiếu trong ảnh tham chiếu I ;

dẫn xuất, tại bộ giải mã, các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng các phép toán dịch chuyển sang phải, trong đó các thông số PROF bên trong bao gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu trong khối video;

thu nhận, tại bộ giải mã, các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên các thông số PROF bên trong; và

thu nhận, tại bộ giải mã, các mẫu dự đoán được tinh chỉnh của khối video dựa trên các mẫu dự đoán ban đầu và các giá trị tinh chỉnh dự đoán.

2. Phương pháp theo điểm 1, trong đó bước dẫn xuất các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng các phép toán dịch chuyển sang phải bao gồm các việc, cho một mẫu trong khối video:

thu nhận, tại bộ giải mã, giá trị gradient ngang dựa trên sự chênh lệch giữa mẫu dự đoán thứ nhất và mẫu dự đoán thứ hai được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận, tại bộ giải mã, giá trị gradient dọc dựa trên sự chênh lệch giữa mẫu dự đoán thứ ba và mẫu dự đoán thứ tư được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận, tại bộ giải mã, các vectơ chuyển động (Motion Vector, MV) điểm điều khiển của một khối mà chứa khối video;

thu nhận, tại bộ giải mã, các thông số mô hình affin dựa trên các MV điểm điều khiển;

thu nhận, tại bộ giải mã, giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc dựa trên các thông số mô hình affin; và

thu nhận, tại bộ giải mã, giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc bởi việc dịch chuyển sang phải theo giá trị dịch chuyển bit, trong đó giá trị dịch chuyển bit bằng 8.

3. Phương pháp theo điểm 2, phương pháp còn bao gồm các bước:

cắt xén, tại bộ giải mã, giá trị chênh lệch chuyển động ngang thành phạm vi đối xứng là $[-31, 31]$; và

cắt xén, tại bộ giải mã, giá trị chênh lệch chuyển động dọc thành phạm vi đối xứng là $[-31, 31]$.

4. Phương pháp theo điểm 2, trong đó bước thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video bao gồm việc:

thu nhận, tại bộ giải mã, các giá trị tinh chỉnh dự đoán dựa trên các giá trị gradient ngang, các giá trị chênh lệch chuyển động ngang, các giá trị gradient dọc, và các giá trị chênh lệch chuyển động dọc.

5. Thiết bị điện toán, bao gồm:

một hoặc nhiều bộ xử lý;

phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp lưu trữ các lệnh có thể thực thi được bởi một hoặc nhiều bộ xử lý, trong đó một hoặc nhiều bộ xử lý được tạo cấu hình để:

thu nhận ảnh tham chiếu I được kết hợp với khối video;

thu nhận các mẫu dự đoán ban đầu của khối video từ khối tham chiếu trong ảnh tham chiếu I;

dẫn xuất các thông số tinh chỉnh dự đoán với luồng quang học (PROF) bên trong của quy trình dẫn xuất PROF bằng cách áp dụng các phép toán dịch chuyển sang phải, trong đó các thông số PROF bên trong bao gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu trong khối video;

thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên các thông số PROF bên trong; và

thu nhận các mẫu dự đoán được tinh chỉnh của khối video dựa trên các mẫu dự đoán ban đầu và các giá trị tinh chỉnh dự đoán.

6. Thiết bị điện toán theo điểm 5, trong đó một hoặc nhiều bộ xử lý được tạo cấu hình để dẫn xuất các thông số PROF bên trong của quy trình dẫn xuất PROF bằng cách áp dụng các phép toán dịch chuyển sang phải còn được tạo cấu hình để, cho một mẫu trong khối video:

thu nhận giá trị gradient ngang dựa trên sự chênh lệch giữa mẫu dự đoán thứ nhất và mẫu dự đoán thứ hai được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận giá trị gradient dọc dựa trên sự chênh lệch giữa mẫu dự đoán thứ ba và mẫu dự đoán thứ tư được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận các vectơ chuyển động (MV) điểm điều khiển của một khối mà chứa khối video;

thu nhận các thông số mô hình afin dựa trên các MV điểm điều khiển;

thu nhận giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc dựa trên các thông số mô hình afin; và

thu nhận giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc bởi việc dịch chuyển sang phải theo giá trị dịch chuyển bit, trong đó giá trị dịch chuyển bit bằng 8.

7. Thiết bị điện toán theo điểm 6, trong đó một hoặc nhiều bộ xử lý còn được tạo cấu hình để:

cắt xén giá trị chênh lệch chuyển động ngang thành phạm vi đối xứng là $[-31, 31]$; và

cắt xén giá trị chênh lệch chuyển động dọc thành phạm vi đối xứng là $[-31, 31]$.

8. Thiết bị điện toán theo điểm 6, trong đó một hoặc nhiều bộ xử lý được tạo cấu hình để thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video còn được tạo cấu hình để:

thu nhận các giá trị tinh chỉnh dự đoán dựa trên các giá trị gradient ngang, các giá trị chênh lệch chuyển động ngang, các giá trị gradient dọc, và các giá trị chênh lệch chuyển động dọc.

9. Phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp lưu trữ luồng bit cần được giải mã bởi phương pháp giải mã tín hiệu video bao gồm các bước:

thu nhận ảnh tham chiếu I được kết hợp với khối video;

thu nhận các mẫu dự đoán ban đầu của khối video từ khối tham chiếu trong ảnh tham chiếu I;

dẫn xuất các thông số tinh chỉnh dự đoán với luồng quang học (PROF) bên trong của quy trình dẫn xuất PROF bằng cách áp dụng các phép toán dịch chuyển sang phải, trong đó các thông số PROF bên trong bao gồm các giá trị gradient ngang, các giá trị gradient dọc, các giá trị chênh lệch chuyển động ngang, và các giá trị chênh lệch chuyển động dọc được dẫn xuất cho các mẫu trong khối video;

thu nhận các giá trị tinh chỉnh dự đoán cho các mẫu trong khối video dựa trên các thông số PROF bên trong; và

thu nhận các mẫu dự đoán được tinh chỉnh của khối video dựa trên các mẫu dự đoán ban đầu và các giá trị tinh chỉnh dự đoán.

10. Phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp theo điểm 9, trong đó phương pháp còn bao gồm các bước, cho một mẫu trong khối video:

thu nhận giá trị gradient ngang dựa trên sự chênh lệch giữa mẫu dự đoán thứ nhất và mẫu dự đoán thứ hai được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận giá trị gradient dọc dựa trên sự chênh lệch giữa mẫu dự đoán thứ ba và mẫu dự đoán thứ tư được kết hợp với một mẫu nói trên trong số các mẫu dự đoán ban đầu;

thu nhận các vectơ chuyển động (MV) điểm điều khiển của một khối mà chứa khối video;

thu nhận các thông số mô hình afin dựa trên các MV điểm điều khiển;

thu nhận giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc dựa trên các thông số mô hình afin; và

thu nhận giá trị chênh lệch chuyển động ngang và giá trị chênh lệch chuyển động dọc bởi việc dịch chuyển sang phải theo giá trị dịch chuyển bit, trong đó giá trị dịch chuyển bit bằng 8.

11. Phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp theo điểm 10, trong đó phương pháp còn bao gồm các bước:

cắt xén giá trị chênh lệch chuyển động ngang thành phạm vi đối xứng là $[-31, 31]$; và

cắt xén giá trị chênh lệch chuyển động dọc thành phạm vi đối xứng là $[-31, 31]$.

12. Phương tiện lưu trữ đọc được bởi máy tính không chuyển tiếp theo điểm 10, trong đó phương pháp còn bao gồm bước:

thu nhận các giá trị tinh chỉnh dự đoán dựa trên các giá trị gradient ngang, các giá trị chênh lệch chuyển động ngang, các giá trị gradient dọc, và các giá trị chênh lệch chuyển động dọc.

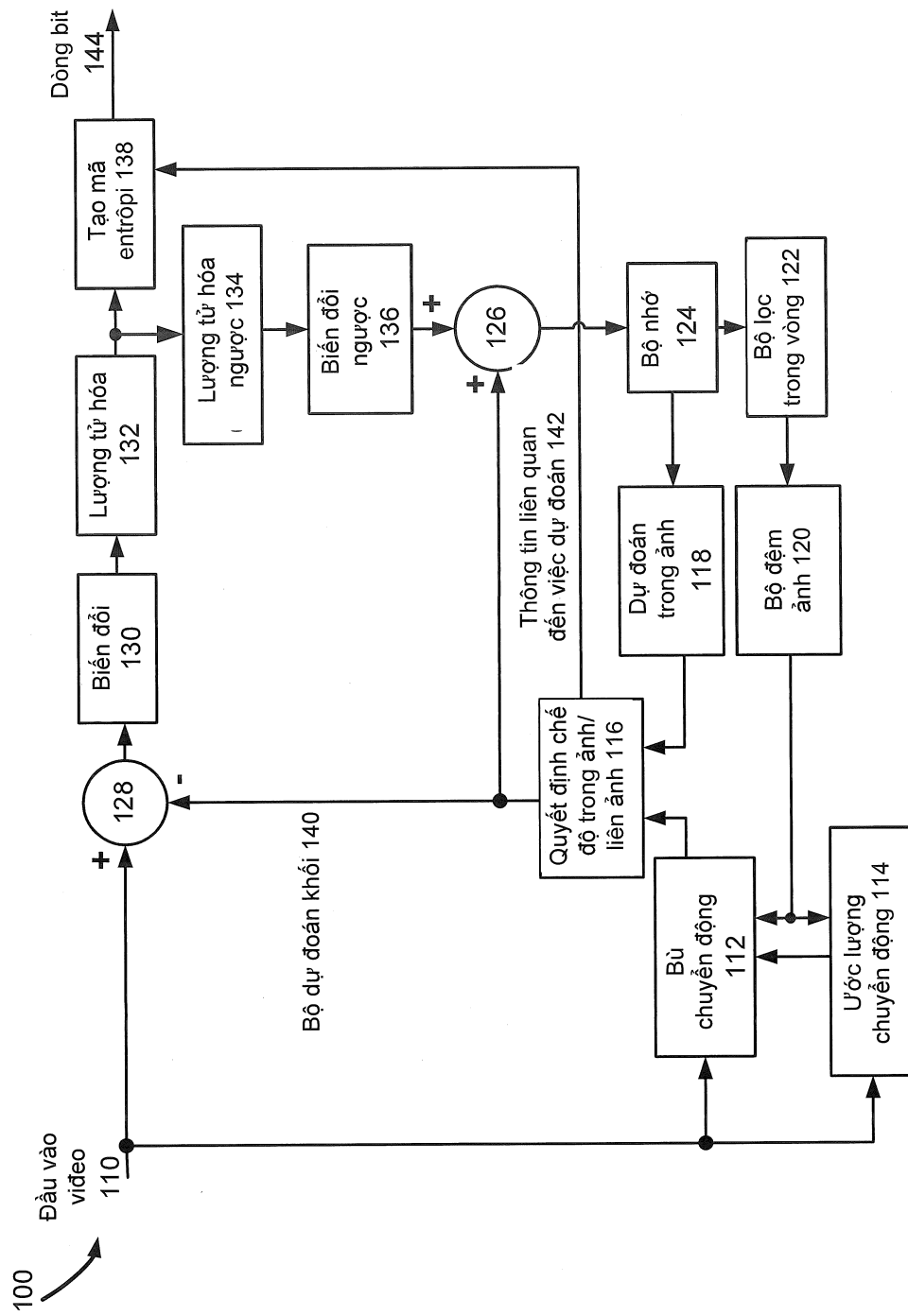
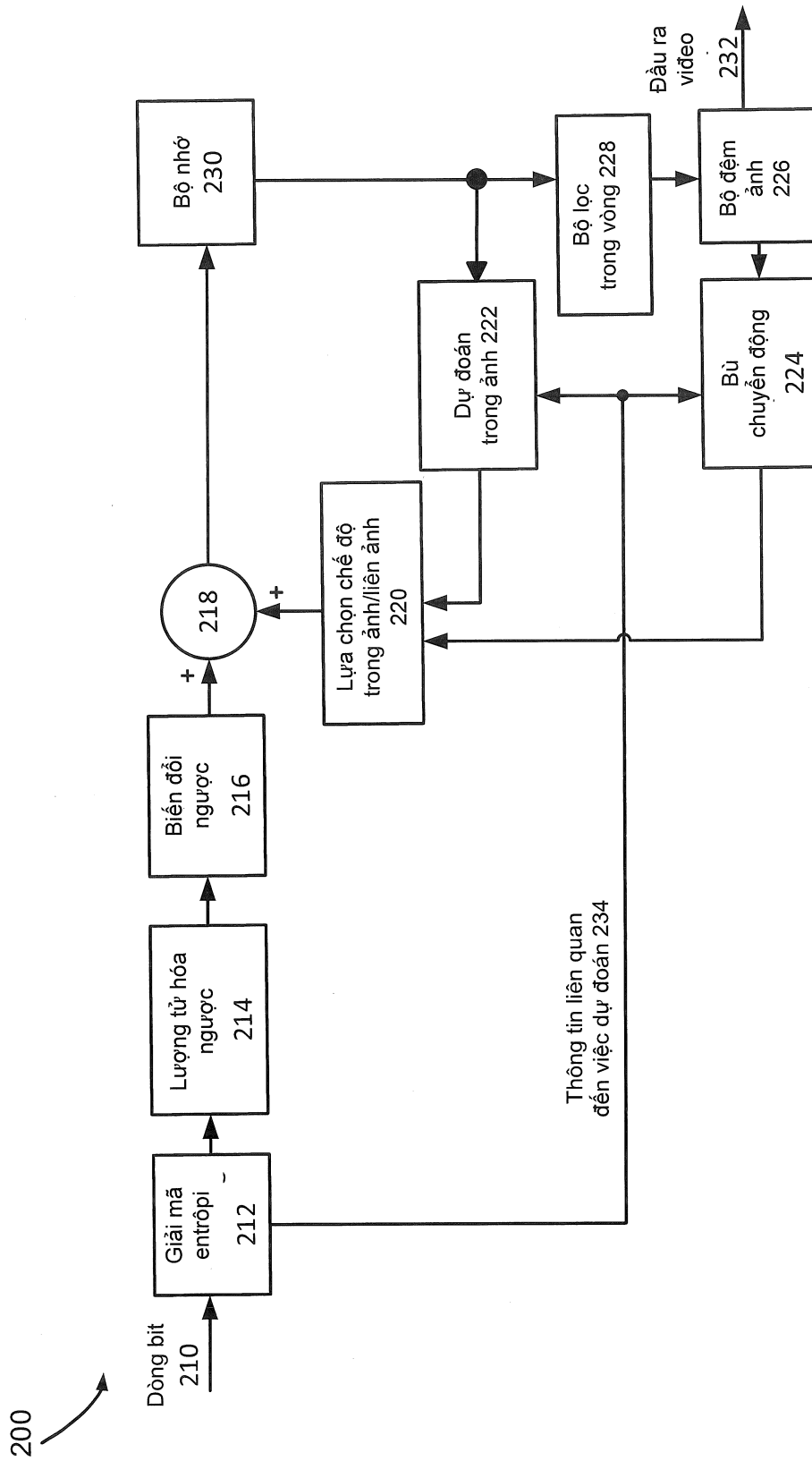


FIG. 1



3 / 14

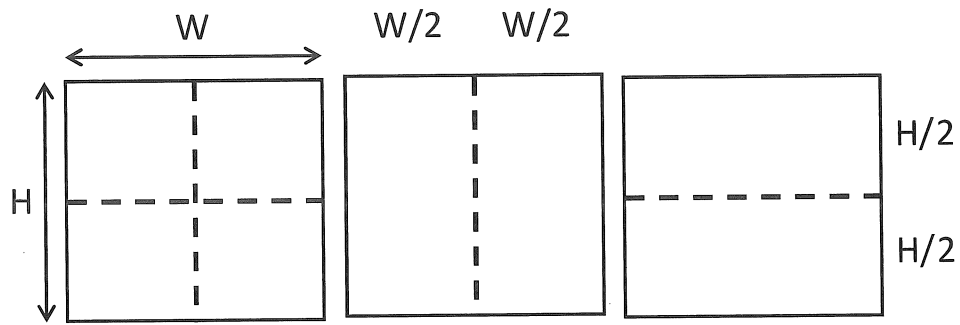


FIG. 3A

FIG. 3B

FIG. 3C

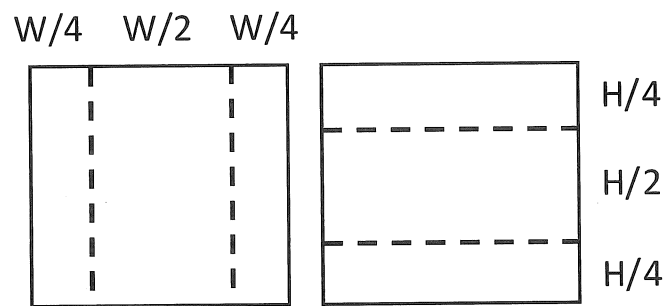


FIG. 3D

FIG. 3E

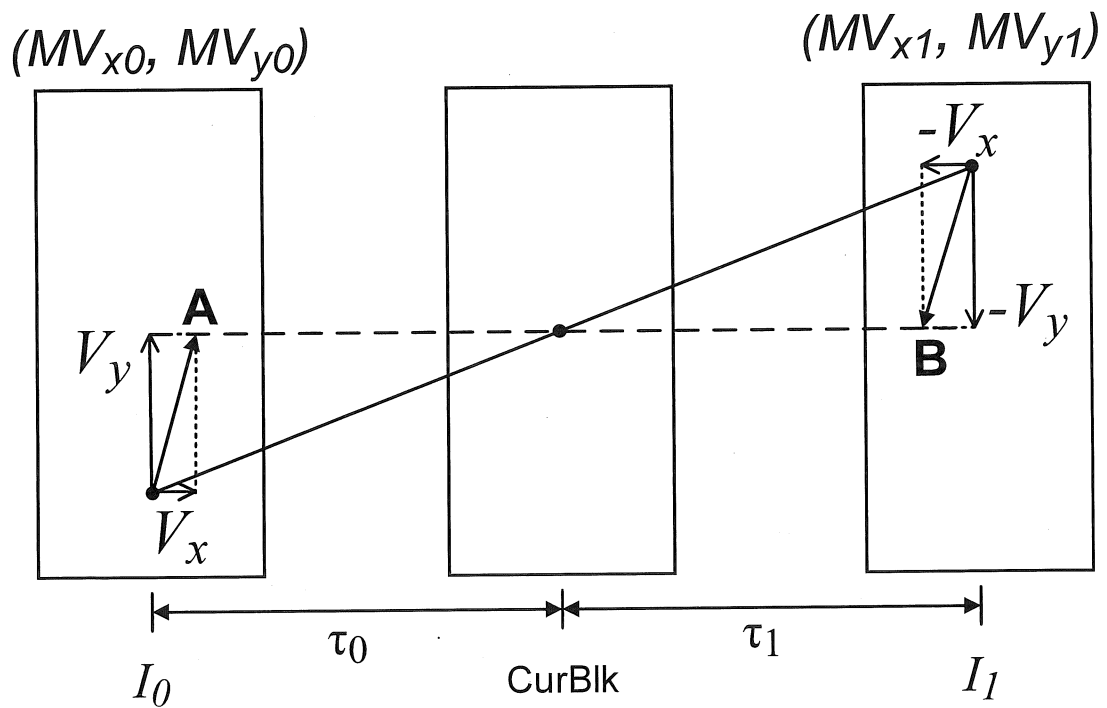


FIG. 4

4 / 14

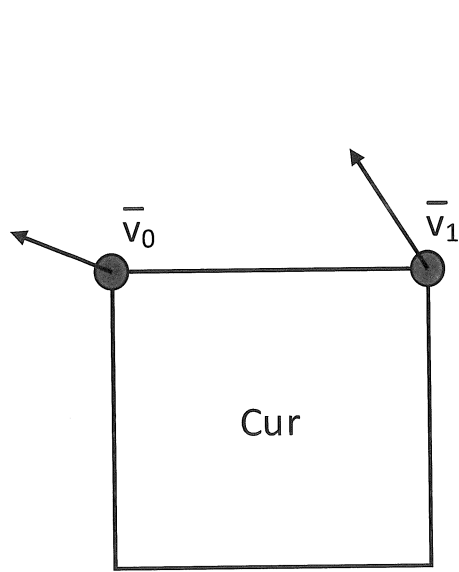


FIG. 5A

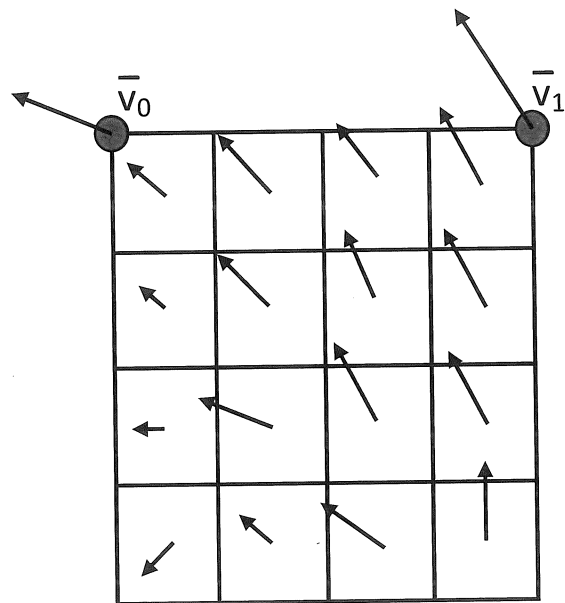


FIG. 5B

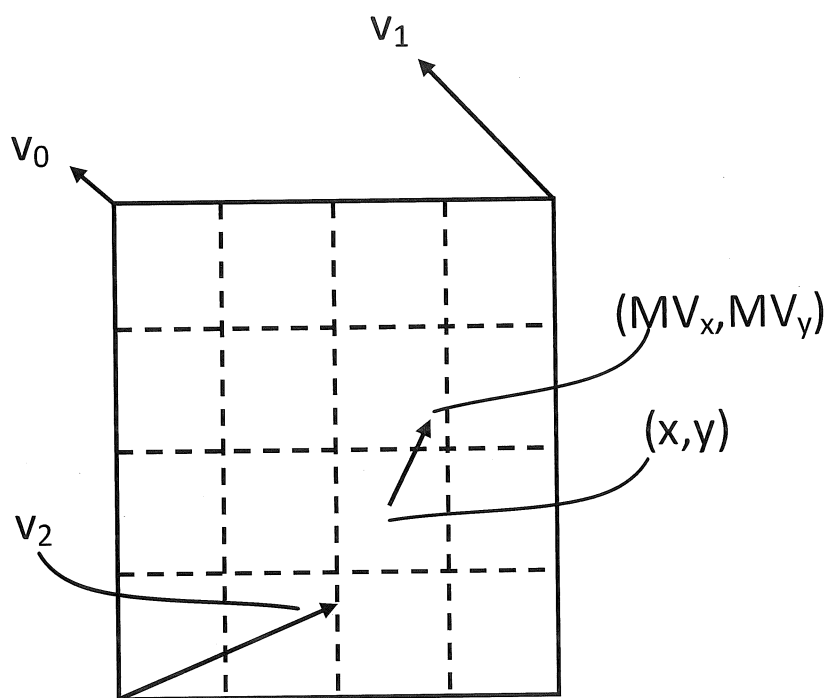


FIG. 6

5 / 14

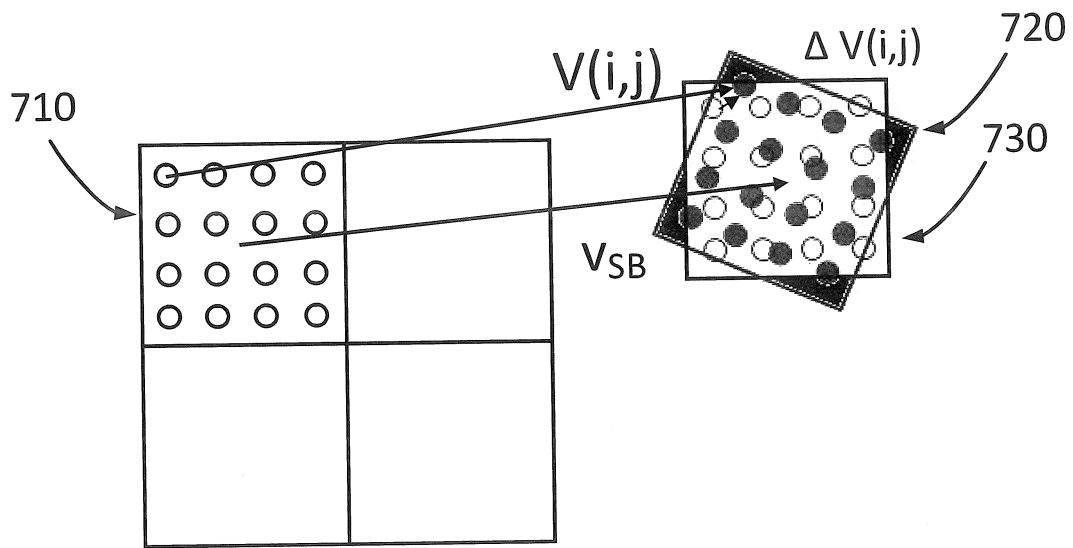


FIG. 7

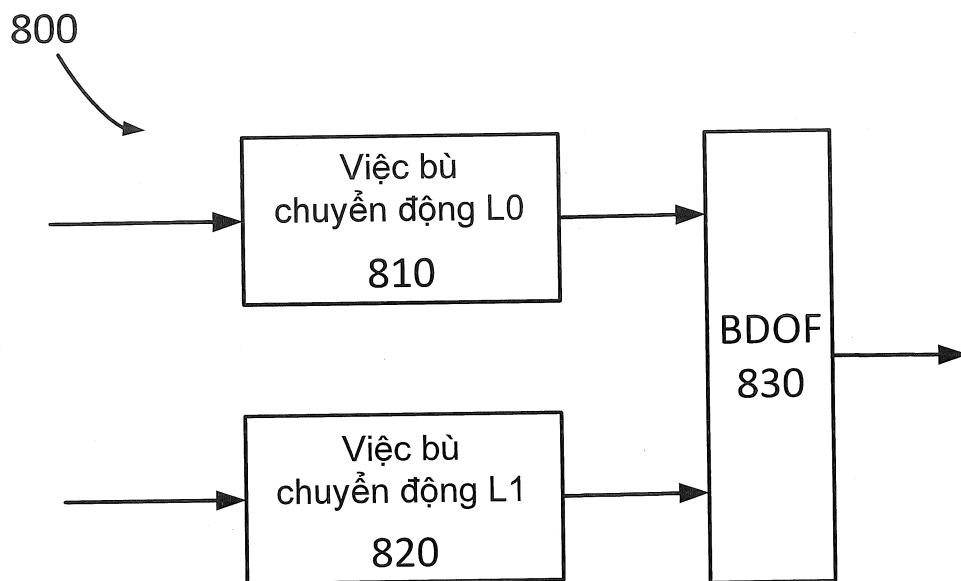


FIG. 8

6 / 14

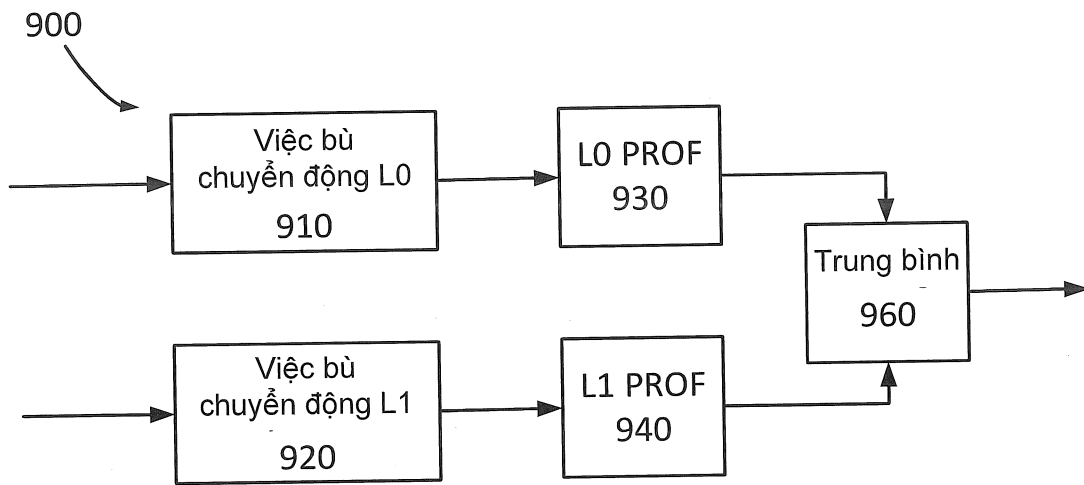


FIG. 9

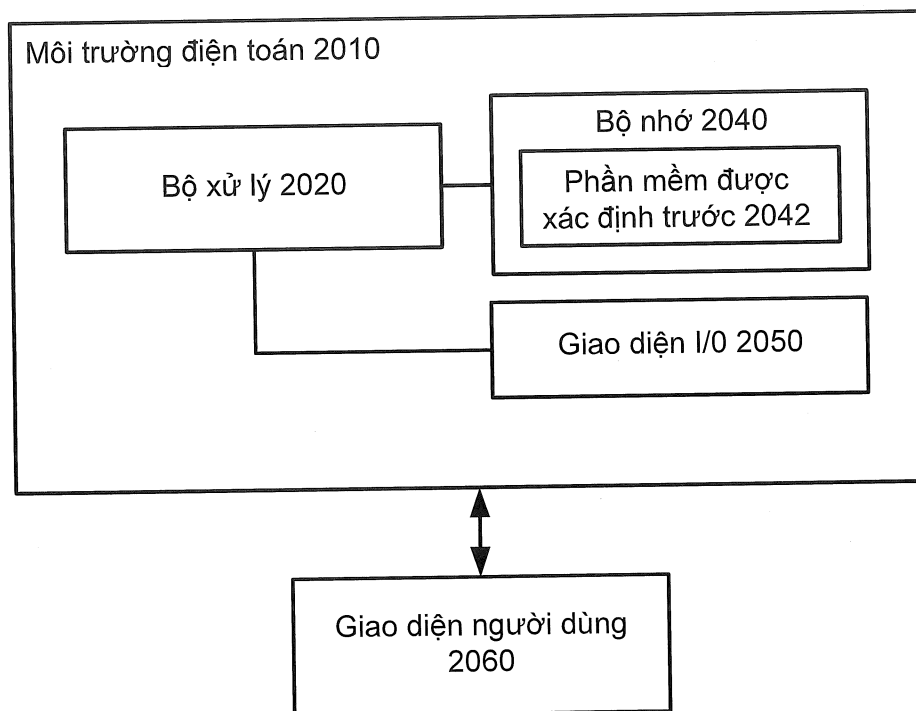


FIG. 20

7 / 14

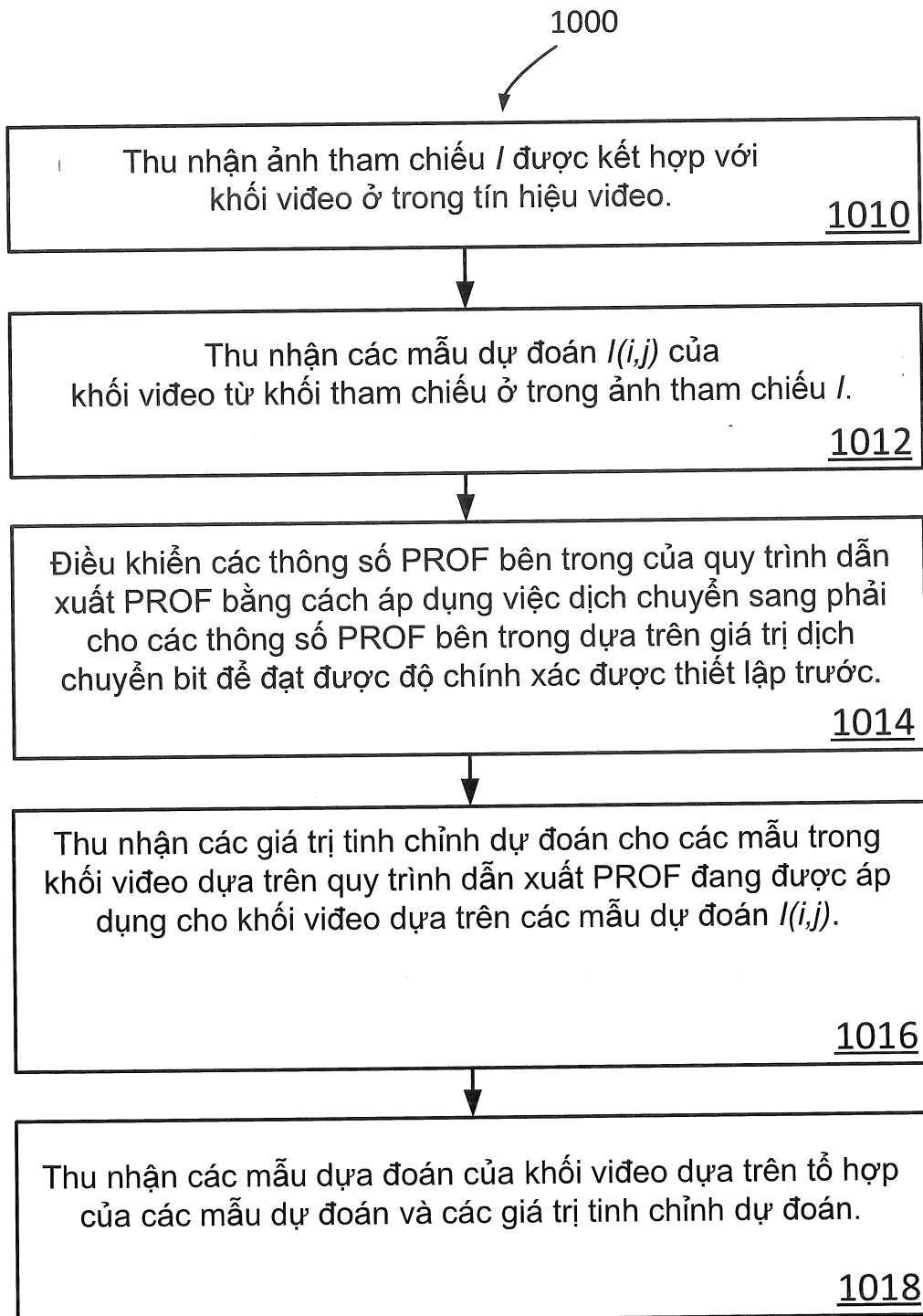


FIG. 10

8 / 14

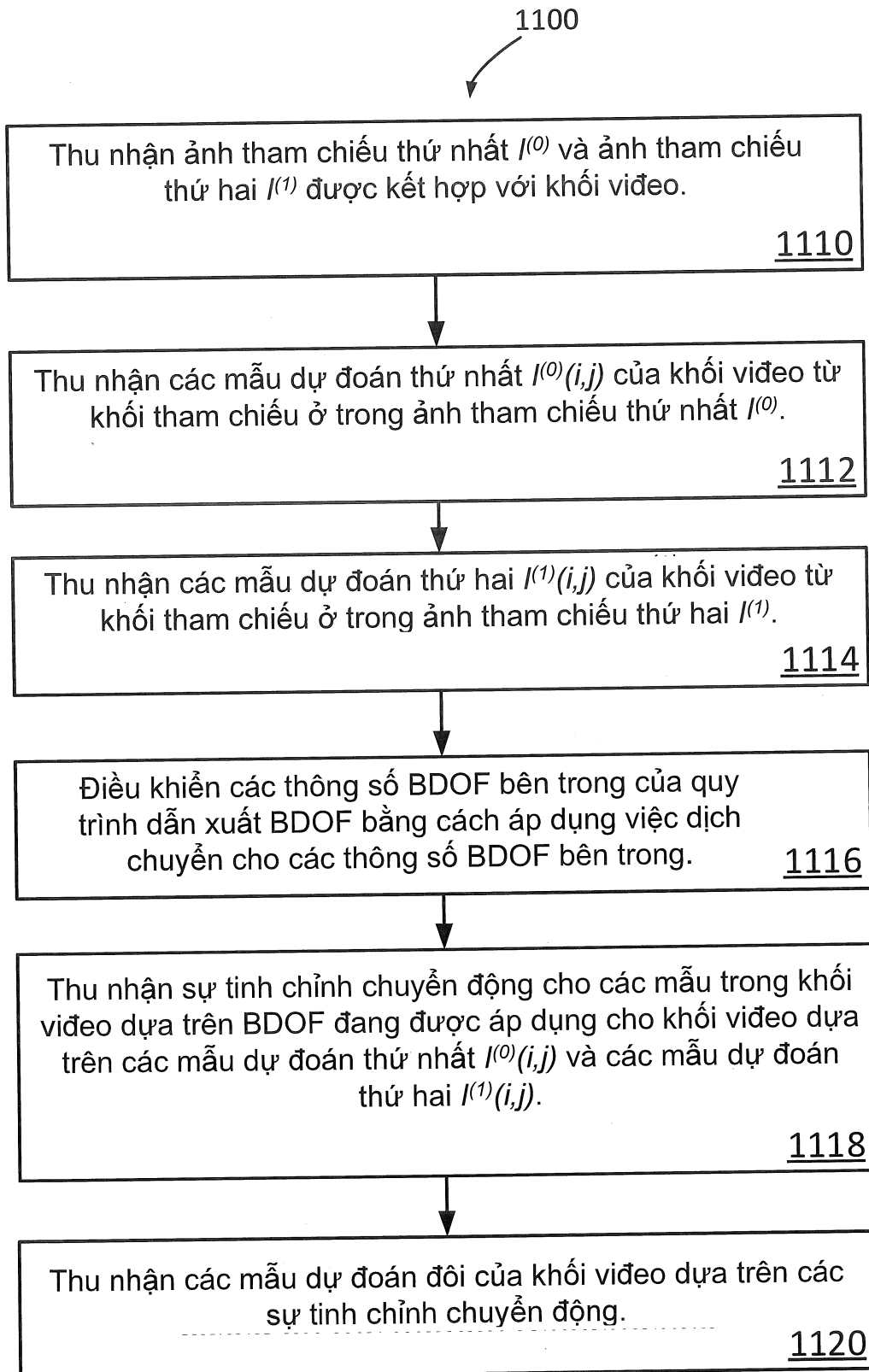


FIG. 11

9 / 14

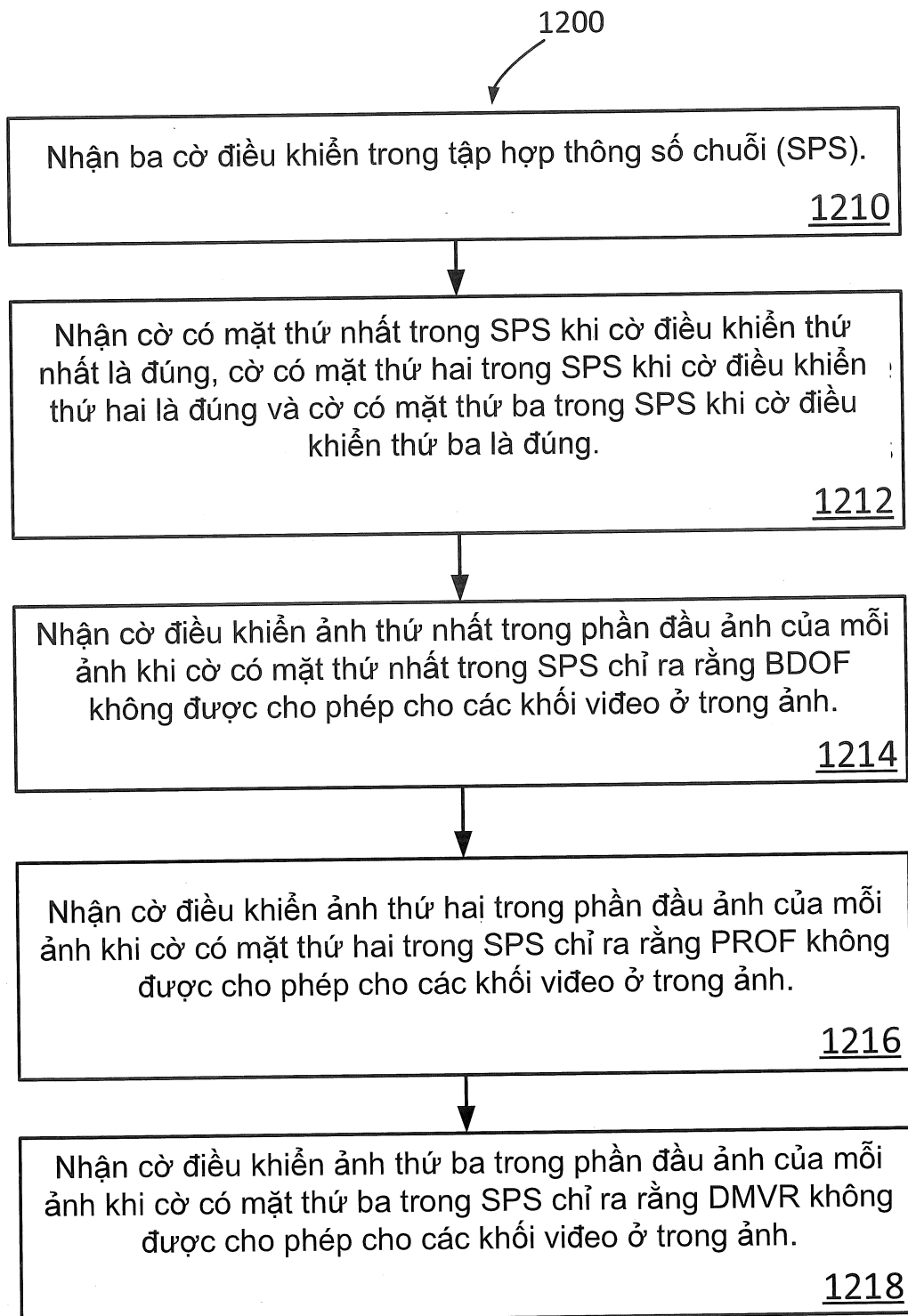


FIG. 12

10 / 14

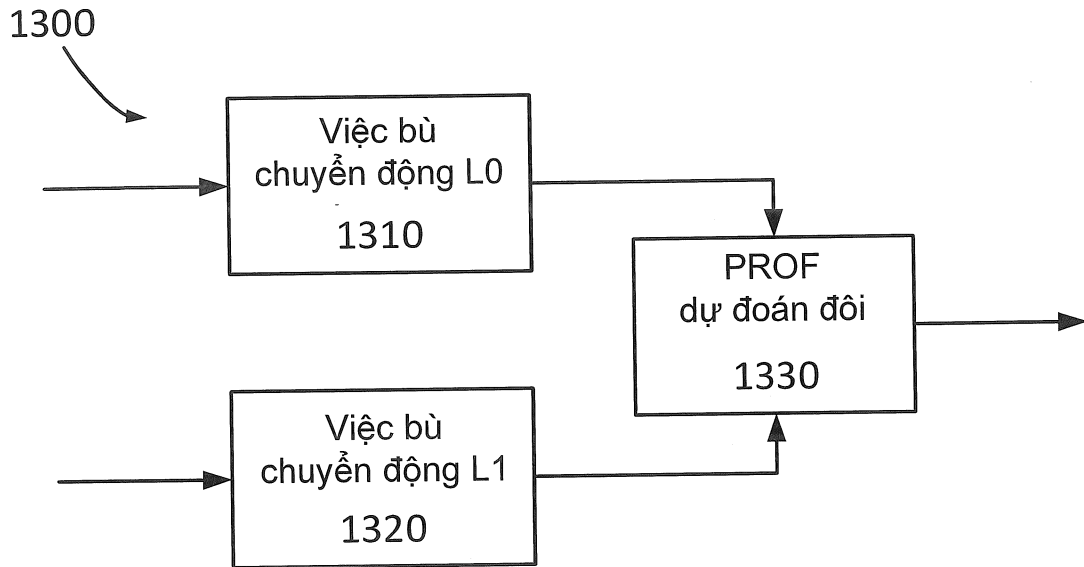


FIG. 13

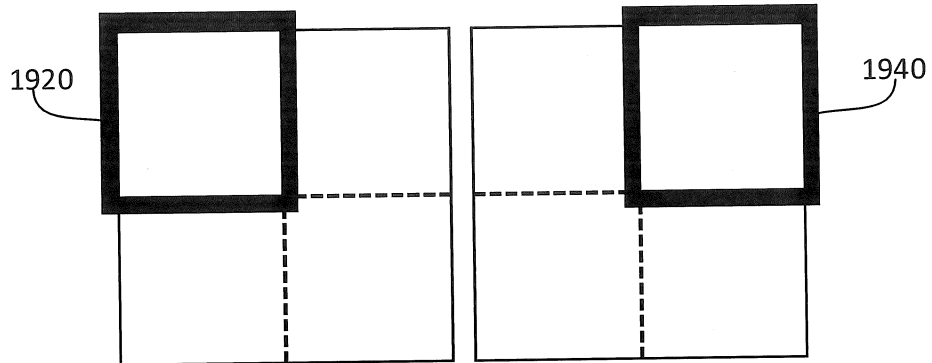
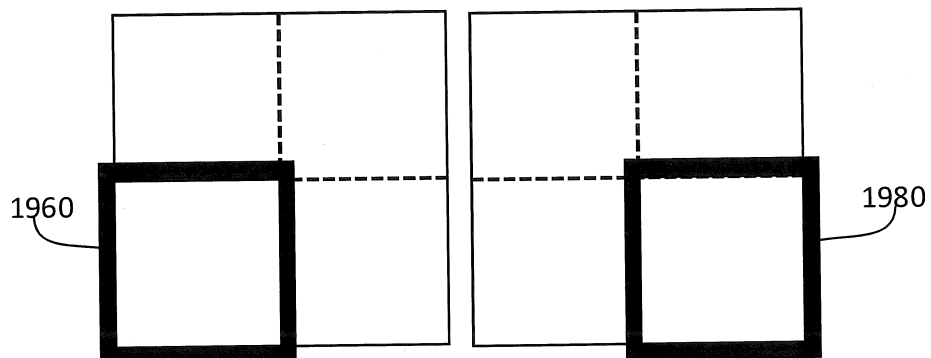


FIG. 19A

FIG. 19B



----- các ranh giới
khối con 4X4

□ các mẫu
được đệm

FIG. 19C

FIG. 19D

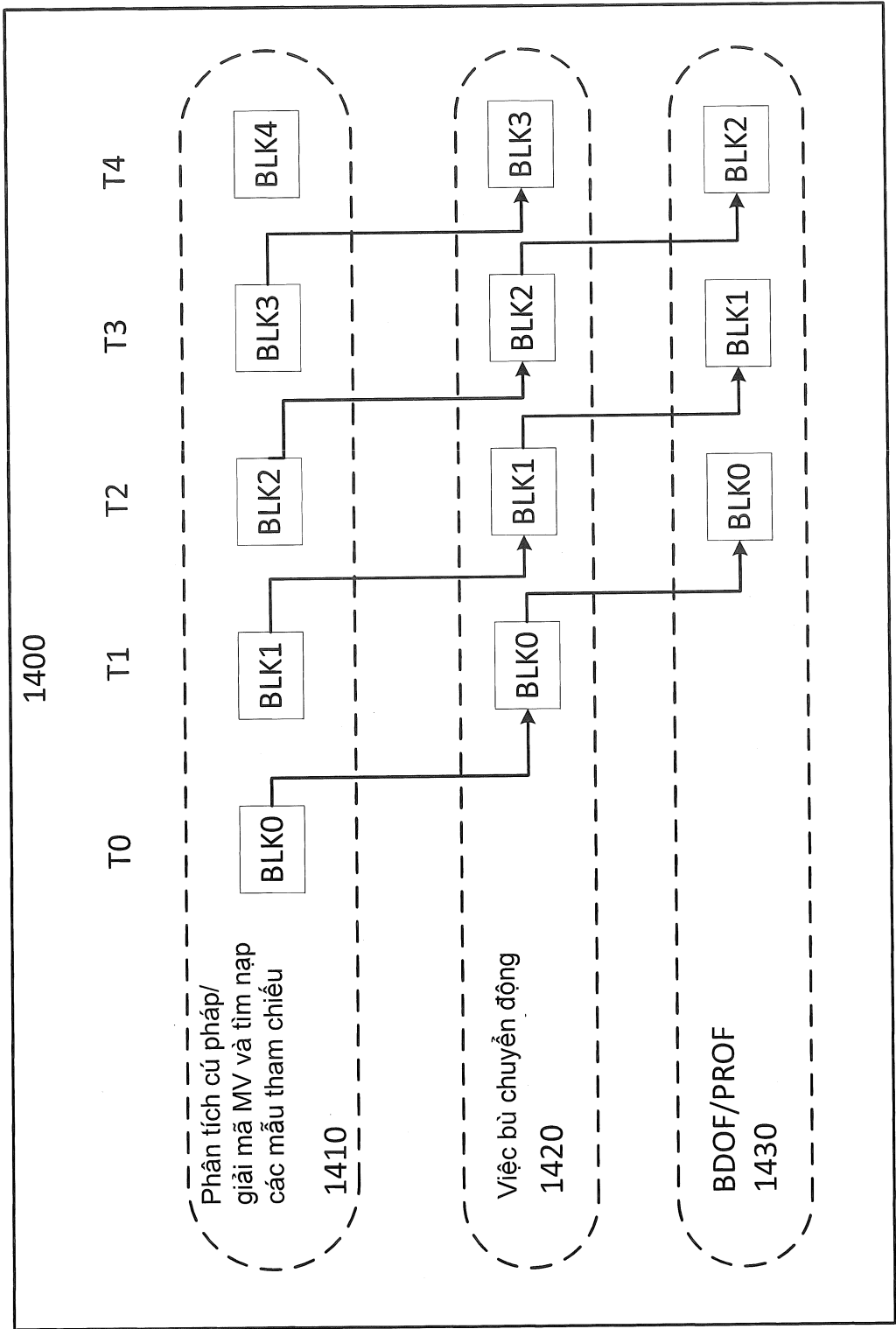


FIG. 14

12 / 14

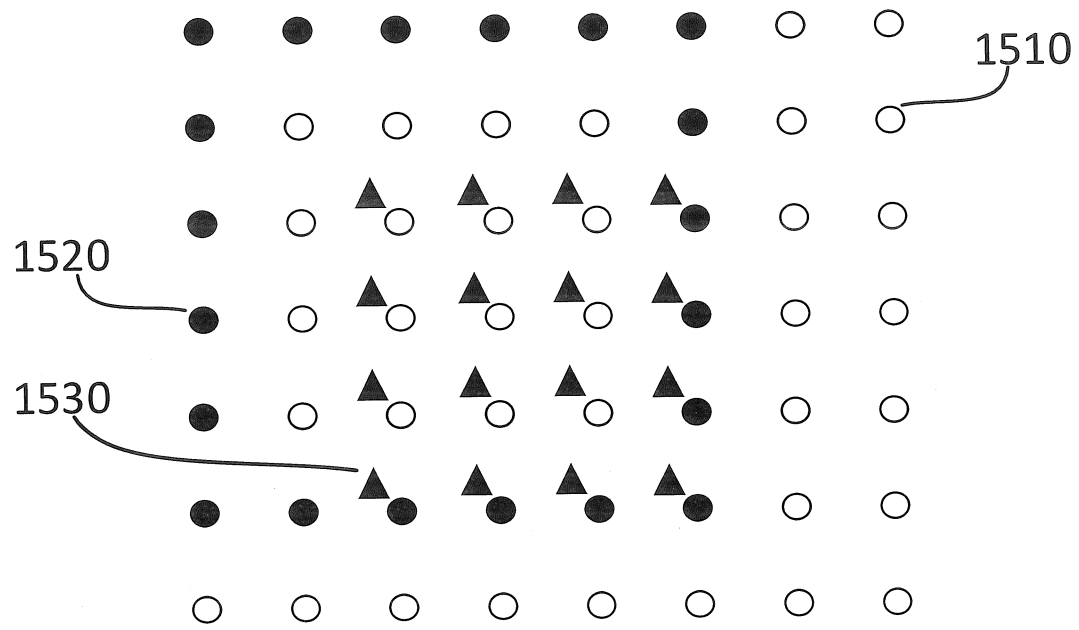


FIG. 15

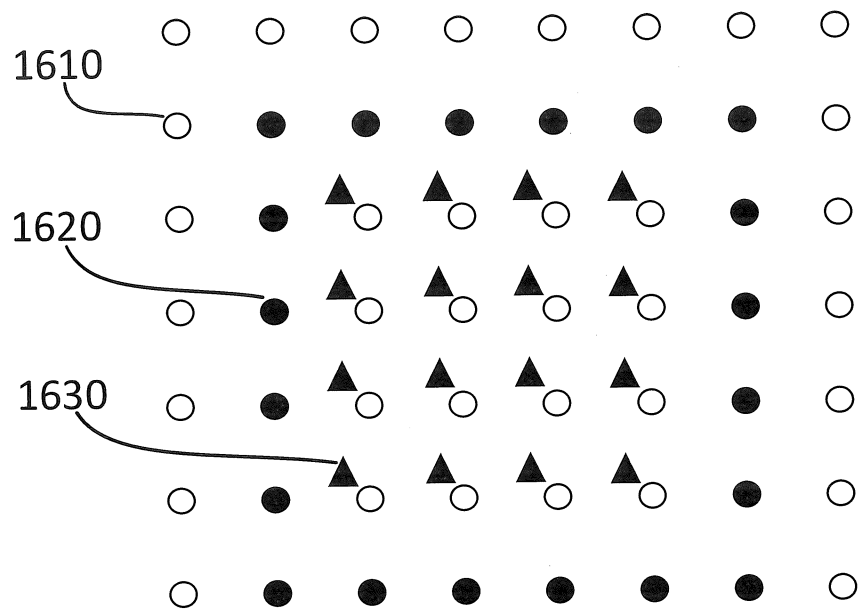


FIG. 16

13 / 14

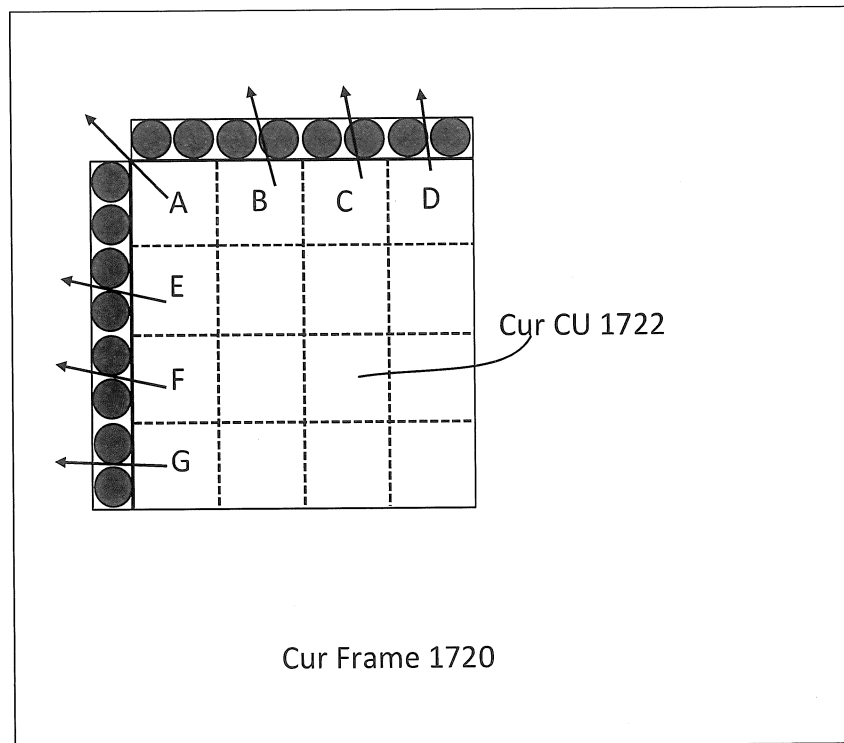


FIG. 17A

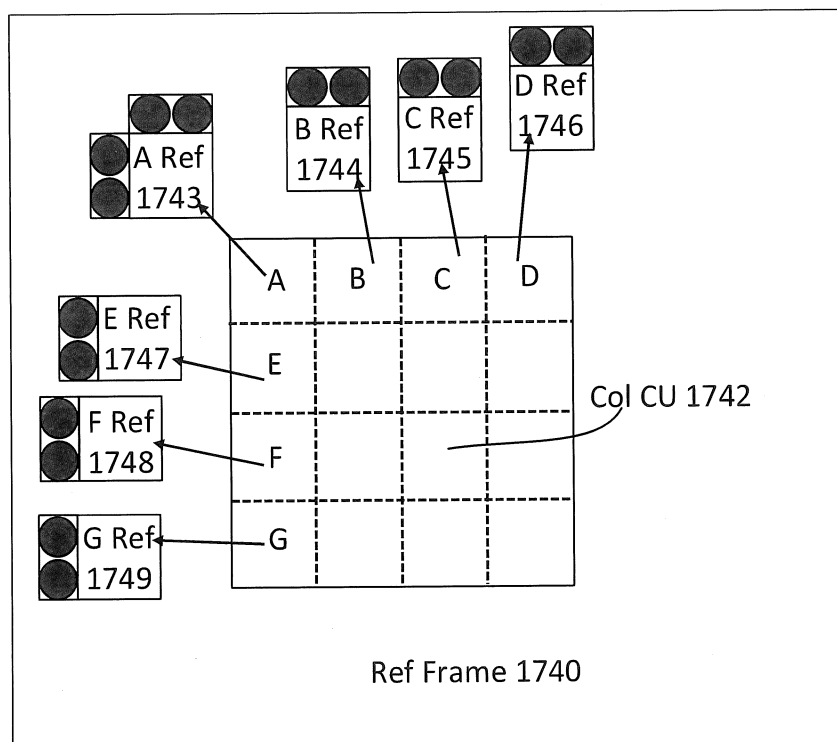


FIG. 17B

14 / 14

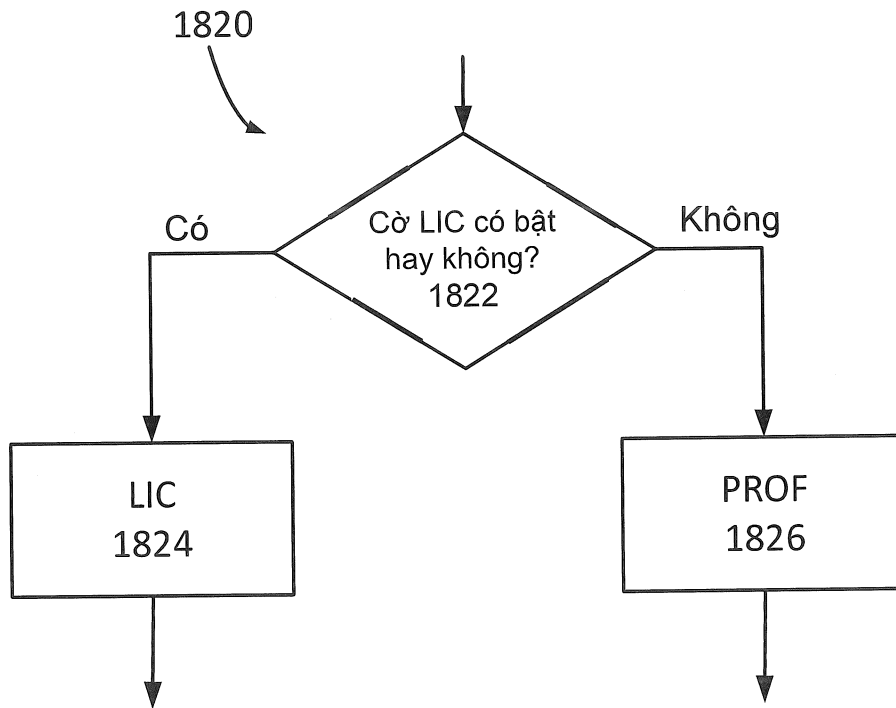


FIG. 18A

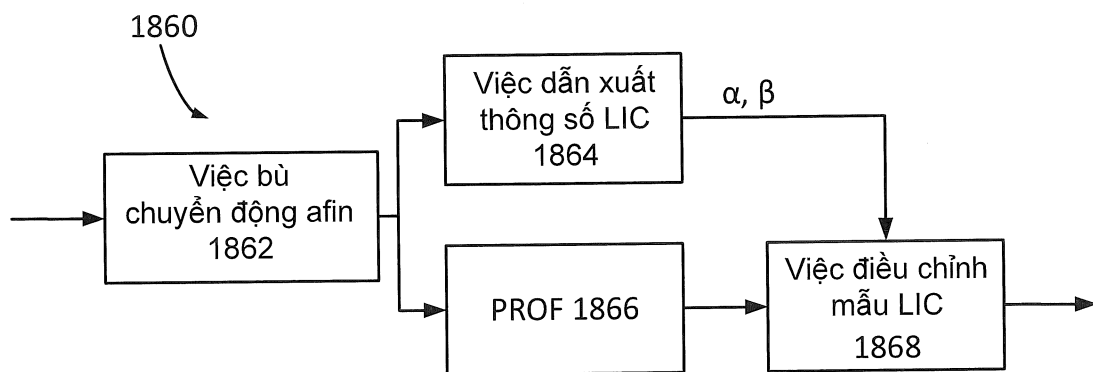


FIG. 18B