



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

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

(11)



1-0048884

(51)^{2020.01} H04W 74/08

(13) B

(21) 1-2022-00297

(22) 03/07/2020

(86) PCT/CN2020/100211 03/07/2020

(87) WO2021/004396 14/01/2021

(30) 201910606607.7 05/07/2019 CN

(45) 25/07/2025 448

(43) 27/06/2022 411A

(73) HUAWEI TECHNOLOGIES CO., LTD. (CN)

Huawei Administration Building, Bantian, Longgang District, Shenzhen, Guangdong
518129, P. R. China

(72) YANG, Mao (CN); LI, Bo (CN); LI, Yunbo (CN).

(74) Công ty Cổ phần Sở hữu công nghiệp INVESTIP (INVESTIP)

(54) PHƯƠNG PHÁP TRUYỀN THÔNG, THỰC THỂ ĐA LIÊN KẾT VÀ PHƯƠNG
TIỆN LƯU TRỮ ĐỌC ĐƯỢC BẰNG MÁY TÍNH

(21) 1-2022-00297

(57) Sáng chế đề xuất phương pháp truyền thông, thực thể đa liên kết (multi-link, ML) và phương tiện lưu trữ đọc được bằng máy tính, và đề cập đến lĩnh vực công nghệ truyền thông, để đảm bảo công bằng trong sự tranh chấp kênh giữa thực thể ML và thực thể liên kết đơn (single-link, SL). Phương pháp truyền thông được áp dụng cho thực thể ML, thực thể ML hỗ trợ liên kết chính và ít nhất một liên kết không chính, bộ đếm chờ truyền được bố trí trên liên kết chính, và không có bộ đếm chờ truyền được bố trí trên liên kết không chính. Phương pháp truyền thông bao gồm: Thực thể ML thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền; và khi giá trị đếm của bộ đếm chờ truyền là 0, thực thể ML gửi đơn vị dữ liệu giao thức lớp vật lý (physical layer protocol data unit, PPDU) trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K liên kết thứ nhất bao gồm liên kết chính và K-1 liên kết không chính thứ nhất, và liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0. Sáng chế áp dụng được cho quá trình truy nhập kênh đa liên kết.

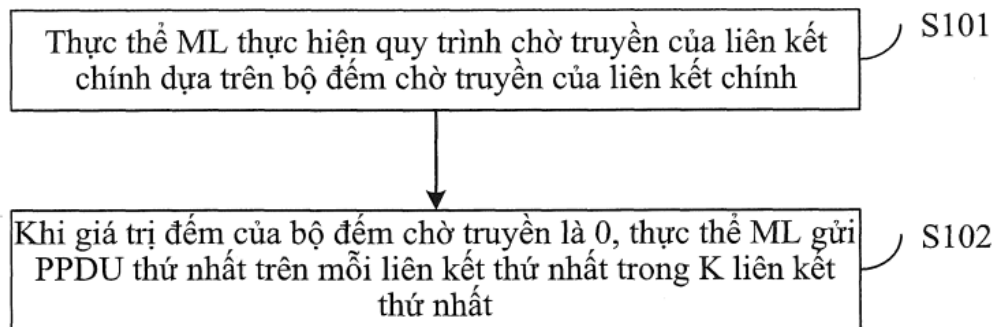


FIG. 5

Lĩnh vực kỹ thuật được đề cập

Sáng chế đề cập đến lĩnh vực công nghệ truyền thông, và cụ thể, là phương pháp truyền thông và bộ máy.

Tình trạng kỹ thuật của sáng chế

Đề đạt được mục tiêu kỹ thuật thông lượng cực cao, tiêu chuẩn của Hội kỹ sư điện và điện tử (institute of electrical and electronics engineers, IEEE) 802.11be bao gồm đa liên kết (multi-link, ML) như là một trong những công nghệ then chốt của nó. Thực thể ML hỗ trợ công nghệ ML có khả năng gửi và nhận qua nhiều dải tần số, sao cho thực thể ML có thể sử dụng băng thông lớn hơn để thực hiện truyền dữ liệu. Điều này cải thiện đáng kể tốc độ thông lượng. Đường dẫn không gian mà thông qua đó thực thể ML thực hiện truyền dữ liệu qua dải tần số có thể được coi là liên kết.

Hiện tại, đối với bất kỳ liên kết trong nhiều liên kết được hỗ trợ bởi thực thể ML, thực thể ML có thể có hai cách truy nhập kênh trên liên kết. Cách 1: Khi giá trị đếm của bộ đếm chờ truyền của liên kết giảm xuống 0, thực thể ML có thể thực hiện truy nhập kênh trên liên kết. Cách 2: Khi bộ đếm chờ truyền của liên kết khác giảm xuống 0, nếu liên kết ở trong trạng thái nghỉ trong PIFS trước, thực thể ML có thể thực hiện truy nhập kênh trên liên kết.

Vì thực thể liên kết đơn (single link, SL) hỗ trợ truyền dữ liệu chỉ trên một liên kết, chỉ khi giá trị đếm của bộ đếm chờ truyền của liên kết được hỗ trợ bởi thực thể SL giảm xuống 0, thực thể SL có thể thực hiện truy nhập kênh trên liên kết.

Vì thế, đối với một liên kết, khả năng mà thực thể ML thu được kênh thông qua sự tranh chấp là lớn hơn khả năng mà thực thể SL thu được kênh thông qua sự tranh chấp. Nói cách khác, khi cả hai thực thể ML và thực thể SL được triển khai trong WLAN, thực thể SL gặp bất lợi trong sự tranh chấp kênh. Điều này ảnh hưởng đến truyền thông thích hợp của thực thể SL.

Bản chất kỹ thuật của sáng chế

Sáng chế đề xuất phương pháp truyền thông và bộ máy, để đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh.

Theo khía cạnh thứ nhất, phương pháp truyền thông được đề xuất. Phương pháp được áp dụng cho thực thể ML, thực thể ML hỗ trợ liên kết chính và ít nhất một liên kết không chính, bộ đếm chờ truyền được bố trí trên liên kết chính, và không có bộ đếm chờ truyền được bố trí trên liên kết không chính. Phương pháp bao gồm: Thực thể ML thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền; và khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, thực thể ML gửi đơn vị dữ liệu giao thức lớp vật lý thứ nhất (physical layer protocol data unit, PPDU) trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K liên kết thứ nhất bao gồm liên kết chính và K-1 liên kết không chính thứ nhất, liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0, và K là số nguyên dương.

Theo giải pháp kỹ thuật trên, vì thực thể ML đặt bộ đếm chờ truyền chỉ trên liên kết chính, khi thực hiện truy nhập kênh, thực thể ML thực hiện quy trình chờ truyền chỉ trên liên kết chính. Theo cách này, thực thể ML không thể thu được kênh thông qua sự tranh chấp trước khi quy trình chờ truyền của liên kết chính kết thúc. Điều này đảm bảo rằng xác suất thu được kênh thông qua sự tranh chấp trên liên kết chính bởi thực thể ML là bằng với xác suất thu được kênh thông qua sự tranh chấp trên liên kết được hỗ trợ của thực thể SL bởi thực thể SL. Vì thế, các giải pháp kỹ thuật được đề xuất trong sáng chế có thể đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Ngoài ra, theo giải pháp kỹ thuật trên, khi liên kết được hỗ trợ của thực thể SL và liên kết chính của thực thể ML là cùng một liên kết, thực thể SL và thực thể ML thực sự thực hiện sự tranh chấp kênh trên cùng một liên kết. Theo cách này, khi thực thể ML thu được thành công kênh thông qua sự tranh chấp trên liên kết chính, thực thể SL không gửi PPDU trên liên kết chính. Điều này đảm bảo sự đồng bộ trong việc nhận và việc gửi được thực hiện bởi thực thể ML trên nhiều liên kết. Ví dụ như, liên kết #1 được sử dụng như liên kết chính. Khi giá trị đếm của bộ đếm chờ truyền của thực thể ML AP trên liên kết #1 là 0, thực thể ML AP gửi PPDU trên liên kết #1 và liên kết #2. Thực thể SL không gửi PPDU đến thực thể ML AP trên liên kết #1. Vì thế, thực thể ML AP có thể nhận một cách đồng bộ các tín hiệu, hoặc gửi một cách đồng bộ các tín hiệu trên liên kết #1 và liên kết #2.

Theo thiết kế khả thi, việc thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ nhất trong K liên kết thứ nhất bao gồm: Thực thể ML gửi PPDU thứ nhất trên kênh khả dụng của mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó kênh khả dụng của liên kết chính bao gồm kênh chính của liên kết chính, và kênh khả dụng của liên kết không chính thứ nhất bao gồm kênh chính của liên kết không chính thứ nhất.

Theo thiết kế khả thi, việc thực thể ML thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền bao gồm: Thực thể ML đợi chu kỳ nghỉ của kênh chính của liên kết chính để đạt tới khoảng thời gian liên khung thứ hai; sau chu kỳ nghỉ của kênh chính của liên kết chính đạt tới khoảng thời gian liên khung thứ hai, mỗi lần kênh chính của liên kết chính ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1; và khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, thực thể ML kết thúc quy trình chờ truyền của liên kết chính.

Theo thiết kế khả thi, việc liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính bao gồm: Kênh chính của liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0.

Theo thiết kế khả thi, kênh chính của liên kết không chính thứ nhất là kênh phụ có tần số thấp nhất, 20 MHz, trong dải tần số tương ứng với liên kết không chính thứ nhất. Ngoài ra, kênh chính của liên kết không chính thứ nhất là kênh phụ có tần số cao nhất, 20 MHz, trong dải tần số tương ứng với liên kết không chính thứ nhất. Nói cách khác, kênh chính của liên kết không chính thứ nhất được tạo cấu hình một cách ngẫu nhiên. Điều này giúp giảm chi phí phát tín hiệu.

Theo thiết kế khả thi, PPDU thứ nhất bao gồm khung điều khiển truy nhập môi trường (media access control, MAC) loại thứ nhất, trong đó khung MAC loại thứ nhất không cần đáp ứng.

Theo thiết kế khả thi, PPDU thứ nhất bao gồm khung MAC loại thứ hai, trong đó khung MAC loại thứ hai cần đáp ứng. Phương pháp còn bao gồm: Thực thể ML nhận khung đáp ứng cho khung MAC loại thứ hai trên một hoặc nhiều liên kết thứ nhất; và nếu một hoặc nhiều liên kết thứ nhất không bao gồm liên kết chính, thực thể ML xác định rằng việc thiết lập của cơ hội truyền (transmission opportunity, TXOP) thất bại;

hoặc nếu một hoặc nhiều liên kết thứ nhất bao gồm liên kết chính, thực thể ML xác định rằng việc thiết lập của TXOP thành công.

Theo thiết kế khả thi, phương pháp còn bao gồm: Thực thể ML xác định N liên kết thứ hai tương ứng với TXOP, trong đó N liên kết thứ hai bao gồm liên kết chính và N-1 liên kết không chính thứ hai, liên kết không chính thứ hai là liên kết không chính thứ nhất mà đáp ứng điều kiện đặt trước, trong đó điều kiện đặt trước bao gồm: Trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất, hoặc trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai, và nhận khung đáp ứng cho khung MAC loại thứ hai. Thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Theo thiết kế khả thi, phương pháp còn bao gồm: Nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết không chính thứ hai, thực thể ML dừng gửi PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thất bại, và tiếp tục gửi, đến khi TXOP kết thúc, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thành công.

Theo thiết kế khả thi, phương pháp còn bao gồm: Nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML dừng gửi PPDU thứ hai trên N liên kết thứ hai; thực thể ML đợi chu kỳ nghỉ của liên kết chính để đạt tới khoảng thời gian liên khung thứ nhất; sau khi chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba, trong đó P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, liên kết không chính thứ ba là liên kết không chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm là thời điểm mà tại đó chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất, và P là số nguyên dương nhỏ hơn hoặc bằng N.

Theo thiết kế khả thi, phương pháp còn bao gồm: Nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML dừng gửi PPDU thứ hai trên N liên kết thứ hai; thực thể ML thực hiện quy trình chờ truyền trên liên kết chính; sau khi quy trình chờ truyền của liên kết chính kết thúc, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba, trong đó P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, liên kết không chính thứ ba là liên kết không

chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính, và P là số nguyên dương nhỏ hơn hoặc bằng N .

Theo khía cạnh thứ hai, phương pháp truyền thông được đề xuất. Phương pháp được áp dụng cho thực thể ML, và thực thể ML hỗ trợ K liên kết thứ nhất. Phương pháp bao gồm: Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2; khi quy trình chờ truyền của liên kết mục tiêu kết thúc, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, liên kết mục tiêu là liên kết thứ nhất mà quy trình chờ truyền kết thúc trước tiên trong K liên kết thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng K ; và nếu sự truyền PPDU thứ nhất thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML không gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ nhất thất bại, hoặc thực thể ML không gửi PPDU thứ hai trên N liên kết thứ hai ở trong chu kỳ thời gian đặt trước.

Theo giải pháp kỹ thuật trên, nếu thực thể ML truyền thất bại PPDU thứ nhất trên một hoặc nhiều liên kết thứ hai, thực thể ML bị cấm gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thất bại; hoặc thực thể ML bị cấm gửi PPDU thứ hai trên N liên kết thứ hai ở trong chu kỳ thời gian đặt trước. Theo cách này, thực thể ML không thể sử dụng nhiều liên kết ở trong chu kỳ thời gian đặt trước. Nếu một liên kết mà thực thể ML không thể sử dụng và ở trong nhiều liên kết được hỗ trợ bởi thực thể SL, ở trong chu kỳ thời gian đặt trước, vì thực thể ML không thể thực hiện sự tranh chấp kênh trên liên kết được hỗ trợ bởi thực thể SL, xác suất mà thực thể SL thu được kênh thông qua sự tranh chấp gia tăng. Điều này đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Theo khía cạnh thứ ba, phương pháp truyền thông được đề xuất. Phương pháp được áp dụng cho thực thể ML, và thực thể ML hỗ trợ K liên kết thứ nhất. Phương pháp bao gồm: Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2; và thực thể

ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà quy trình chờ truyền đã kết thúc và ở trong trạng thái nghỉ trong khoảng thời gian liên khung trước thời điểm thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng M .

Theo giải pháp kỹ thuật trên, mặc dù ML thực hiện quy trình chờ truyền trên tất cả K liên kết thứ nhất, liên kết thứ hai được sử dụng để gửi PPDU thứ nhất cần thỏa mãn điều kiện mà quy trình chờ truyền của liên kết thứ hai đã kết thúc. Nói cách khác, trên một liên kết, thực thể ML có thể thu được kênh thông qua sự tranh chấp chỉ sau khi thực thể ML kết thúc quy trình chờ truyền trên liên kết. So sánh với công nghệ thông thường trong đó thực thể ML có thể thu được kênh thông qua sự tranh chấp trên một liên kết kể cả nếu quy trình chờ truyền trên liên kết chưa kết thúc, giải pháp kỹ thuật trong sáng chế giảm xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp trên một liên kết. Điều này đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Theo thiết kế khả thi, thời điểm thứ nhất là thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, và liên kết mục tiêu là liên kết thứ hai mà quy trình chờ truyền kết thúc cuối cùng trong N liên kết thứ hai.

Theo khía cạnh thứ tư, phương pháp truyền thông được đề xuất. Phương pháp được áp dụng cho thực thể ML, và thực thể ML hỗ trợ K liên kết thứ nhất. Phương pháp bao gồm: Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2; và khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0, hoặc tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ hai trước thời điểm hiện tại, và N là số nguyên dương nhỏ hơn hoặc bằng M .

Theo giải pháp kỹ thuật trên, mặc dù thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, thực thể ML có thể thu được một cách thành công kênh thông qua sự tranh chấp chỉ khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0 hoặc tổng các giá trị đếm

của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0. Nói cách khác, đối với thực thể ML, các giá trị đếm của các bộ đếm chờ truyền của một hoặc nhiều liên kết thứ nhất cần phải nhỏ hơn 0. Điều này yêu cầu một hoặc nhiều liên kết thứ nhất nghỉ trong thời gian tương đối dài. Theo cách này, xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp bị giảm. Việc xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp bị giảm làm suy yếu lợi thế của thực thể ML so với thực thể SL trong sự tranh chấp kênh, đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Theo thiết kế khả thi, việc thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất bao gồm: đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, thực thể ML đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung; và sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất đi 1.

Theo thiết kế khả thi, giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất bao gồm các số nguyên âm.

Theo thiết kế khả thi, phương pháp còn bao gồm: đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm mục tiêu đi 1, trong đó bộ đếm mục tiêu được tạo cấu hình để ghi lại tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất. Theo cách này, thực thể ML có thể nhận biết một cách trực tiếp tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất bằng cách sử dụng bộ đếm mục tiêu.

Theo khía cạnh thứ năm, phương pháp truyền thông được đề xuất. Phương pháp được áp dụng cho thực thể ML, thực thể ML hỗ trợ nhiều liên kết, và nhiều liên kết mỗi liên kết đóng vai trò là liên kết thứ nhất đến lượt theo thứ tự tuần hoàn đặt trước. Phương pháp bao gồm: Thực thể ML thực hiện quy trình chờ truyền trên liên kết thứ nhất; và sau khi quy trình chờ truyền của liên kết thứ nhất kết thúc, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó N liên kết thứ hai bao

gồm liên kết thứ nhất và N-1 liên kết khả dụng, liên kết khả dụng ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết thứ nhất, và N là số nguyên dương.

Theo giải pháp kỹ thuật trên, mỗi lần truy nhập kênh được thực hiện, thực thể ML thực hiện quy trình chờ truyền chỉ trên liên kết thứ nhất. Nói cách khác, thực thể ML thực hiện sự tranh chấp kênh trên chỉ một liên kết duy nhất. Xác suất thực thể ML thu được kênh thông qua sự tranh chấp trên một liên kết bằng với xác suất thực thể SL thu được kênh thông qua sự tranh chấp trên một liên kết. Theo cách này, công bằng cho thực thể SL trong sự tranh chấp kênh được đảm bảo, và vì thế truyền thông thích hợp của thực thể SL được đảm bảo.

Theo khía cạnh thứ sáu, thực thể ML được đề xuất. Thực thể ML có thể bao gồm môđun được tạo cấu hình để thực hiện phương pháp/thao tác/bước/hành động được mô tả trong bất kỳ thiết kế theo khía cạnh thứ nhất đến khía cạnh thứ năm trong sự tương ứng một-một. Môđun trên có thể là mạch điện phần cứng, hoặc phần mềm, hoặc có thể được triển khai bằng cách sử dụng mạch điện phần cứng kết hợp với phần mềm.

Theo khía cạnh thứ bảy, thực thể ML được đề xuất. Thực thể ML bao gồm bộ xử lý và bộ truyền nhận, và bộ xử lý được tạo cấu hình để thực hiện thao tác xử lý theo phương pháp truyền thông theo bất kỳ thiết kế theo khía cạnh thứ nhất đến khía cạnh thứ năm. Bộ truyền nhận được tạo cấu hình để được điều khiển bởi bộ xử lý để thực hiện thao tác gửi và nhận theo phương pháp truyền thông theo bất kỳ thiết kế theo khía cạnh thứ nhất đến khía cạnh thứ năm.

Theo khía cạnh thứ tám, phương tiện lưu trữ đọc được bằng máy tính được đề xuất. Phương tiện lưu trữ đọc được bằng máy tính được tạo cấu hình để lưu trữ các lệnh. Khi các lệnh được đọc bởi máy tính, máy tính được tạo cấu hình để thực hiện phương pháp truyền thông theo bất kỳ thiết kế theo khía cạnh thứ nhất đến khía cạnh thứ năm.

Theo khía cạnh thứ chín, sản phẩm chương trình máy tính được đề xuất. Sản phẩm chương trình máy tính bao gồm các lệnh. Khi máy tính đọc các lệnh, máy tính thực hiện phương pháp truyền thông theo bất kỳ thiết kế khả thi theo khía cạnh thứ nhất đến khía cạnh thứ tư.

Theo khía cạnh thứ mười, chip được đề xuất. Chip bao gồm mạch điện xử lý và chân truyền nhận. Chip hỗ trợ liên kết chính và ít nhất một liên kết không chính. Bộ đếm

chờ truyền được bố trí trên liên kết chính, và không có bộ đếm chờ truyền được bố trí trên liên kết không chính. Mạch điện xử lý được tạo cấu hình để thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền. Chân truyền nhận được tạo cấu hình để: khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, gửi PPDU thứ nhất trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K liên kết thứ nhất bao gồm liên kết chính và K-1 liên kết không chính thứ nhất, liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0, và K là số nguyên dương.

Theo khía cạnh thứ mười một, chip được đề xuất. Chip bao gồm mạch điện xử lý và chân truyền nhận. Chip hỗ trợ K liên kết thứ nhất. Mạch điện xử lý được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Chân truyền nhận được tạo cấu hình để: khi quy trình chờ truyền của liên kết mục tiêu kết thúc, gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, liên kết mục tiêu là liên kết thứ nhất mà quy trình chờ truyền kết thúc trước tiên trong K liên kết thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng K. Chân truyền nhận còn được tạo cấu hình để: nếu sự truyền PPDU thứ nhất thất bại trên một hoặc nhiều liên kết thứ hai, bỏ qua gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ nhất thất bại, hoặc bỏ qua gửi PPDU thứ hai trên N liên kết thứ hai ở trong chu kỳ thời gian đặt trước.

Theo khía cạnh thứ mười hai, chip được đề xuất. Chip bao gồm mạch điện xử lý và chân truyền nhận. Chip hỗ trợ K liên kết thứ nhất. Mạch điện xử lý được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Chân truyền nhận được tạo cấu hình để gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà quy trình chờ truyền đã kết thúc và ở trong trạng thái nghỉ trong khoảng thời gian liên khung trước thời điểm thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng M.

Theo khía cạnh thứ mười ba, chip được đề xuất. Chip bao gồm mạch điện xử lý

và chân truyền nhận. Chip hỗ trợ K liên kết thứ nhất. Mạch điện xử lý được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Chân truyền nhận được tạo cấu hình để: khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0, hoặc tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0, gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ hai trước thời điểm hiện tại, và N là số nguyên dương nhỏ hơn hoặc bằng M.

Theo khía cạnh thứ mười bốn, chip được đề xuất. Chip bao gồm mạch điện xử lý và chân truyền nhận. Chip hỗ trợ nhiều liên kết, và nhiều liên kết mỗi liên kết đóng vai trò là liên kết thứ nhất đến lượt theo thứ tự tuần hoàn đặt trước. Mạch điện xử lý được tạo cấu hình để thực hiện quy trình chờ truyền trên liên kết thứ nhất. Chân truyền nhận được tạo cấu hình để: sau khi quy trình chờ truyền của liên kết thứ nhất kết thúc, gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó N liên kết thứ hai bao gồm liên kết thứ nhất và N-1 liên kết khả dụng, liên kết khả dụng ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết thứ nhất, và N là số nguyên dương.

Đối với các hiệu quả kỹ thuật được đem lại bởi bất kỳ thiết kế theo khía cạnh thứ sáu đến khía cạnh thứ mười bốn, tham chiếu đến các hiệu quả có lợi theo phương pháp tương ứng được đề xuất trên. Các chi tiết không được mô tả lần nữa ở đây.

Mô tả văn tắt các hình vẽ

Fig.1 là sơ đồ dạng giản đồ của quy trình chờ truyền theo phương án của sáng chế;

Fig.2 là sơ đồ dạng giản đồ của cấu trúc khung của PPDU theo phương án của sáng chế;

Fig.3 là sơ đồ dạng giản đồ của kịch bản truyền thông ML theo phương án của sáng chế;

Fig.4 là sơ đồ dạng giản đồ của kịch bản truyền thông ML khác theo phương án của sáng chế;

Fig.5 là lưu đồ của phương pháp truyền thông theo phương án của sáng chế;

Fig.6 là lưu đồ của phương pháp truyền thông khác theo phương án của sáng chế;

Fig.7 là sơ đồ dạng giản đồ của văn kịch bản truyền thông ML khác theo phương án của sáng chế;

Fig.8 là sơ đồ dạng giản đồ của văn kịch bản truyền thông ML khác theo phương án của sáng chế;

Fig.9(a) là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.9(b) là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.10 là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.11(a) là lưu đồ của phương pháp truyền thông khác theo phương án của sáng chế;

Fig.11(b) là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.12 là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.13 là lưu đồ của văn phương pháp truyền thông khác theo phương án của sáng chế;

Fig.14 là sơ đồ dạng giản đồ của cấu trúc của thực thể ML theo phương án của sáng chế; và

Fig.15 là sơ đồ dạng giản đồ của cấu trúc của thực thể ML theo phương án của sáng chế.

Mô tả chi tiết bản sáng chế

Theo các sự mô tả của sáng chế, trừ khi được chỉ định theo cách khác “/” nghĩa là “hoặc”. Ví dụ như, A/B có thể biểu diễn A hoặc B. “Và/hoặc” trong bản mô tả mô tả mối quan hệ liên hợp cho việc mô tả các đối tượng được liên hợp và biểu diễn rằng có thể có ba mối quan hệ. Ví dụ như, A và/hoặc B có thể biểu diễn ba trường hợp sau: Chỉ có A tồn tại, cả hai A và B tồn tại, và chỉ có B tồn tại. Ngoài ra, “ít nhất một” nghĩa là một hoặc nhiều, và “nhiều” nghĩa là hai hoặc nhiều. Các thuật ngữ như “thứ nhất” và “thứ hai” không giới hạn số lượng hoặc trình tự thực thi, và các thuật ngữ như “thứ nhất”

và "thứ hai" không chỉ ra sự khác biệt rõ ràng.

Cần lưu ý rằng, trong sáng chế, các thuật ngữ như là “ví dụ” hoặc “ví dụ như” được sử dụng để biểu diễn việc đưa ra ví dụ, minh họa, hoặc mô tả. Bất kỳ phương án hoặc sơ đồ thiết kế được mô tả như “ví dụ” hoặc “ví dụ như” trong sáng chế không nên được giải thích là được ưu tiên hơn hoặc có nhiều lợi thế hơn so với phương án hoặc sơ đồ thiết kế khác. Cụ thể, việc sử dụng từ “ví dụ” và “ví dụ như” nhằm để trình bày khái niệm tương đối theo cách cụ thể.

Để dễ dàng hiểu, phần sau trước tiên mô tả vắn tắt một số thuật ngữ công nghệ theo các phương án của sáng chế.

1. Tập dịch vụ cơ bản (basic service set, BSS)

BSS được sử dụng để mô tả nhóm các thiết bị mà có thể giao tiếp với nhau trong mạng cục bộ không dây (wireless local area network, WLAN). WLAN có thể bao gồm nhiều BSS. Mỗi BSS có mã định danh duy nhất. Mã định danh duy nhất được gọi là mã định danh tập dịch vụ cơ bản (BSSID).

Một BSS có thể bao gồm nhiều trạm (station, STA). Trạm có thể là điểm truy nhập (access point, AP) hoặc trạm điểm không truy nhập (non-access point station, non-AP STA). Một cách tùy chọn, một BSS có thể bao gồm một AP và nhiều non-AP STA được liên hợp với AP.

AP cũng được gọi là điểm truy nhập không dây hoặc điểm phát sóng. AP có thể là bộ định tuyến không dây, bộ truyền nhận không dây, bộ chuyển mạch không dây, hoặc tương tự.

Non-AP STA có thể có các tên khác nhau như là đơn vị thuê bao, thiết bị đầu cuối truy nhập, trạm di động, thiết bị di động, thiết bị đầu cuối, và thiết bị người dùng. Trong ứng dụng thực tế, non-AP STA có thể là điện thoại tế bào, điện thoại thông minh, đường dây thuê bao vô tuyến (wireless local loop, WLL), và thiết bị cầm tay khác hoặc thiết bị máy tính mà có chức năng truyền thông mạng cục bộ không dây.

2. Cơ chế chờ truyền

Tiêu chuẩn IEEE 802.11 cho phép nhiều người dùng chia sẻ cùng một phương tiện truyền. Bộ truyền kiểm tra tính khả dụng của phương tiện truyền trước khi gửi dữ liệu. Tiêu chuẩn IEEE 802.11 sử dụng đa truy nhập có tránh va chạm (carrier sense multiple access with collision avoidance, CSMA/CA) để triển khai sự tranh chấp kênh.

Để tránh sự va chạm, CSMA/CA theo cơ chế chờ truyền.

Cơ chế chờ truyền trên kênh đơn được mô tả dưới đây. Trước khi thiết bị gửi bản tin, thiết bị có thể lựa chọn số ngẫu nhiên từ 0 đến cửa sổ tranh chấp (contention window, CW), và sử dụng số ngẫu nhiên như giá trị khởi tạo của bộ đếm chờ truyền. Sau khi chu kỳ nghỉ của kênh đạt tới khoảng thời liên khung truyền (arbitration inter-frame space, AIFS), giá trị đếm của bộ đếm chờ truyền giảm đi 1 mỗi lần kênh nghỉ trong khe thời gian (timeslot). Trước khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, nếu kênh bận trong một khe thời gian, bộ đếm chờ truyền dừng đếm. Sau đó, nếu kênh thay đổi từ bận sang nghỉ về trạng thái và chu kỳ nghỉ của kênh đạt tới AIFS, bộ đếm chờ truyền lại tiếp tục đếm. Khi giá trị đếm của bộ đếm chờ truyền là 0, quy trình chờ truyền kết thúc, và thiết bị có thể bắt đầu truyền dữ liệu.

Ví dụ được đề xuất cho sự mô tả với tham chiếu đến Fig.1, giả sử rằng giá trị khởi tạo của bộ đếm chờ truyền là 5, và sau khi chu kỳ nghỉ của kênh đạt tới AIFS, bộ đếm chờ truyền bắt đầu thực hiện chờ truyền. Mỗi lần kênh ở trong trạng thái nghỉ trong một khe thời gian, giá trị đếm của bộ đếm chờ truyền giảm đi 1 đến khi giá trị đếm của bộ đếm chờ truyền là 0. Sau khi giá trị đếm của bộ đếm chờ truyền là 0, thiết bị thu được một cách thành công kênh thông qua sự tranh chấp, và thiết bị có thể gửi PPDU trên kênh.

3. PPDU

Fig.2 là sơ đồ dạng giản đồ của cấu trúc khung của PPDU trong tiêu chuẩn 802.11ax. PPDU bao gồm trường huấn luyện ngắn kế thừa (legacy-short training field, L-STF), trường huấn luyện dài kế thừa (legacy-long training field, L-LTF), trường tín hiệu kế thừa (legacy-signal field, L-SIG), trường phát tín hiệu kế thừa lặp lại (repeated legacy-signal field, RL-SIG), trường tín hiệu hiệu quả cao A (high efficiency-signal field A, HE-SIG A), trường tín hiệu hiệu quả cao B (high efficiency-signal field B, HE-SIG B), trường huấn luyện ngắn hiệu quả cao (high efficiency-short training field, HE-STF), và trường huấn luyện dài hiệu quả cao (high efficiency-long training field, HE-LTF), trường dữ liệu (data), và trường mở rộng gói (packet extension, PE).

4. TXOP

TXOP là đơn vị cơ bản trong truy nhập kênh không dây. TXOP bao gồm thời điểm ban đầu và thời lượng tối đa giới hạn TXOP. Ở trong giới hạn TXOP, trạm mà thu

được TXOP có thể không thực hiện sự tranh chấp kênh lặp lại, và liên tục sử dụng kênh để truyền nhiều khung dữ liệu.

5. Cơ chế yêu cầu để gửi (request to send, RTS)/sẵn sàng để gửi (clear to send, CTS)

Cơ chế RTS/CTS được sử dụng để giải quyết vấn đề các trạm ẩn, để tránh xung đột tín hiệu giữa nhiều trạm.

Trước khi gửi khung dữ liệu, đầu truyền trước tiên gửi khung RTS theo cách quảng bá, để chỉ ra rằng đầu truyền là để gửi khung dữ liệu đến đầu nhận chỉ định ở trong thời lượng chỉ định. Sau khi nhận khung RTS, đầu nhận gửi khung CTS theo cách quảng bá để xác nhận sự truyền được thực hiện bởi đầu truyền. Trạm khác mà nhận khung RTS hoặc khung CTS không gửi khung vô tuyến đến khi thời lượng chỉ định kết thúc.

6. Thực thể ML

Thực thể ML có khả năng gửi và nhận qua nhiều dải tần số. Ví dụ như, nhiều dải tần số bao gồm nhưng không bị giới hạn ở dải tần số 2,4 GHz, dải tần số 5 GHz, và dải tần số 6 GHz. Đường dẫn không gian mà thông qua đó thực thể ML thực hiện truyền dữ liệu qua dải tần số có thể được coi là liên kết. Nói cách khác, thực thể ML hỗ trợ truyền thông đa liên kết.

Cần hiểu rằng, đối với thực thể ML, mỗi liên kết được hỗ trợ bởi thực thể ML tương ứng với một dải tần số.

Thực thể ML có thể cũng được gọi là thực thể ML STA. Thực thể ML bao gồm nhiều STA. Nhiều STA trong thực thể ML có thể có cùng một địa chỉ MAC, hoặc có thể có các địa chỉ MAC khác nhau. Nhiều STA trong thực thể ML có thể được định vị tại cùng một vị trí vật lý, hoặc có thể được định vị tại các vị trí vật lý khác nhau.

Mỗi STA trong thực thể ML có thể thiết lập liên kết cho truyền thông. Như được thể hiện trên Fig.3, thực thể ML A bao gồm trạm A1 đến trạm AN, và thực thể ML B bao gồm trạm B1 đến trạm BN. Trạm A1 giao tiếp với trạm B1 thông qua liên kết 1, trạm A2 giao tiếp với trạm B2 thông qua liên kết 2, và bằng cách tương tự, trạm AN giao tiếp với trạm BN thông qua liên kết N.

Khi khoảng cách tần số giữa nhiều dải tần số được hỗ trợ bởi thực thể ML là nhỏ, việc gửi tín hiệu qua một dải tần số bởi thực thể ML ảnh hưởng nghiêm trọng đến việc

nhận tín hiệu qua dải tần số khác. Vì thế, để đảm bảo truyền thông thích hợp, khi thực thể ML thực hiện truyền thông qua nhiều liên kết tại cùng một thời gian, thực thể ML cần phải nhận các tín hiệu trên nhiều liên kết tại cùng một thời gian, hoặc gửi các tín hiệu trên nhiều liên kết tại cùng một thời gian. Trên Fig.4, thực thể ML A gửi các PPDU trên liên kết thứ nhất và liên kết thứ hai tại cùng một thời gian. Sau đó, thực thể ML A nhận các khung xác nhận khối (block acknowledgement, block ack, BA) phản hồi lại bởi thực thể ML B trên liên kết thứ nhất và liên kết thứ hai tại cùng một thời gian.

Theo phương án này của sáng chế, các PPDU được gửi bởi thực thể ML trên các liên kết khác nhau có thể là giống nhau hoặc khác nhau.

Theo phương án này của sáng chế, khi thực thể ML thực hiện truyền thông qua nhiều liên kết tại cùng một thời gian, các mã định danh lưu lượng (traffic identifier, TID) tương ứng với nhiều liên kết có thể là giống nhau hoặc khác nhau.

Nếu các STA trong thực thể ML là các AP, thực thể ML có thể được gọi là thực thể ML AP. Nếu các STA trong thực thể ML là các non-AP STA, thực thể ML có thể được gọi là thực thể ML non-AP STA hoặc thực thể ML non-AP. Theo phương án này của sáng chế, trừ khi được chỉ định theo cách khác, thực thể ML có thể là thực thể ML AP, hoặc có thể là thực thể ML non-AP.

Non-AP STA mà ở trong thực thể ML non-AP và ở trên liên kết có thể được liên hợp với AP mà ở trong thực thể ML AP và ở trên cùng một liên kết, sao cho non-AP STA trong thực thể ML non-AP trên liên kết có thể giao tiếp với AP trong thực thể ML AP trên cùng một liên kết.

Cần hiểu rằng mối quan hệ liên hợp có thể được thiết lập giữa thực thể ML AP và thực thể ML non-AP, để đảm bảo truyền thông thích hợp giữa thực thể ML AP và thực thể ML non-AP.

Cần lưu ý rằng mối quan hệ liên hợp giữa thực thể ML AP và thực thể ML non-AP bao gồm: mối quan hệ liên hợp giữa trạm trong thực thể ML AP trên liên kết và trạm trong thực thể ML non-AP trên cùng một liên kết.

Sự triển khai của việc thiết lập mối quan hệ liên hợp giữa thực thể ML non-AP và thực thể ML AP không bị giới hạn theo phương án này của sáng chế. Ví dụ như, thực thể ML non-AP và thực thể ML AP thiết lập trên một liên kết mối quan hệ liên hợp của liên kết. Ngoài ra, thực thể ML non-AP và thực thể ML AP thiết lập trên một liên kết

mối quan hệ liên hợp của nhiều liên kết giữa thực thể ML non-AP và thực thể ML AP.

Đối với sự triển khai cụ thể của việc thiết lập mối quan hệ liên hợp của một liên kết giữa thực thể ML non-AP và thực thể ML AP, tham chiếu đến sự triển khai của việc thiết lập mối quan hệ liên hợp giữa AP và non-AP STA trong công nghệ thông thường. Các chi tiết không được mô tả lần nữa ở đây.

7. Thực thể SL

Thực thể SL là STA mà hỗ trợ chỉ một liên kết duy nhất. Thực thể SL có thể là STA kế thừa (legacy), nói cách khác, STA mà hỗ trợ chỉ tiêu chuẩn 802.11 hiện có nhưng không hỗ trợ tiêu chuẩn 802.11 thế hệ tiếp theo.

Phần trên mô tả vắn tắt các thuật ngữ kỹ thuật liên quan đến sáng chế, và các chi tiết không được mô tả lần nữa trong phần sau.

Các giải pháp kỹ thuật trong sáng chế được áp dụng cho WLAN. Tiêu chuẩn được sử dụng cho WLAN có thể là tiêu chuẩn IEEE 802.11, ví dụ như, tiêu chuẩn 802.11ax hoặc tiêu chuẩn 802.11 thế hệ tiếp theo. Các kịch bản mà các giải pháp kỹ thuật trong sáng chế áp dụng được bao gồm: kịch bản truyền thông giữa các thực thể ML và kịch bản truyền thông giữa thực thể ML và thực thể SL.

Ví dụ như, kịch bản truyền thông giữa các thực thể ML có thể là kịch bản truyền thông giữa thực thể ML non-AP và thực thể ML AP, kịch bản truyền thông giữa các thực thể ML non-AP, hoặc kịch bản truyền thông giữa các thực thể ML AP.

Ví dụ như, kịch bản truyền thông giữa thực thể ML và thực thể SL có thể là kịch bản truyền thông giữa thực thể ML non-AP và AP kế thừa, kịch bản truyền thông giữa thực thể ML không phải AP và non-AP STA kế thừa, kịch bản truyền thông giữa thực thể ML AP và AP kế thừa, hoặc kịch bản truyền thông giữa thực thể ML non-AP và non-AP STA kế thừa.

Phần sau mô tả cụ thể các giải pháp kỹ thuật được đề xuất theo các phương án của sáng chế với tham chiếu đến các hình vẽ đi kèm của bản mô tả.

Fig.5 thể hiện phương pháp truyền thông theo phương án của sáng chế. Phương pháp bao gồm các bước sau.

S101. Thực thể ML thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền của liên kết chính.

Thực thể ML hỗ trợ liên kết chính và ít nhất một liên kết không chính. Bộ đếm

chờ truyền được bố trí trên liên kết chính, và không có bộ đếm chờ truyền được bố trí trên liên kết không chính. Cần hiểu rằng, vì bộ đếm chờ truyền được bố trí chỉ trên liên kết chính, thực thể ML có thể thực hiện quy trình chờ truyền chỉ trên liên kết chính.

Cần hiểu rằng, nếu thực thể ML hỗ trợ liên kết chính và ít nhất một liên kết không chính, thực thể ML hỗ trợ hai cách truy nhập kênh. Một là cách truy nhập kênh liên kết đơn, nói cách khác, thực thể ML thực hiện truy nhập kênh chỉ trên liên kết chính. Cách khác là cách truy nhập kênh đa liên kết, nói cách khác, thực thể ML thực hiện truy nhập kênh trên liên kết chính và liên kết không chính. Trong sáng chế thực tế, thực thể ML có thể lựa chọn cách truy nhập kênh dựa trên trạng thái kênh, công suất, tải trọng dịch vụ, và tương tự. Ví dụ như, để tiết kiệm năng lượng, thực thể ML sử dụng cách truy nhập kênh liên kết đơn. Ngoài ra, để cải thiện thông lượng, thực thể ML sử dụng cách truy nhập kênh đa liên kết. Phương án này của sáng chế mô tả chủ yếu cách truy nhập kênh đa liên kết.

Một cách tùy chọn, liên kết chính của thực thể ML có thể được tạo cấu hình một cách rõ ràng. Cần hiểu rằng việc tạo cấu hình một cách rõ ràng liên kết chính của thực thể ML là linh hoạt.

Ví dụ như, thực thể AP có thể gửi thông tin chỉ dẫn đến thực thể ML non-AP được liên hợp với thực thể ML AP, để chỉ ra thông tin về liên kết chính. Thông tin về liên kết chính có thể bao gồm mã định danh/chỉ số của liên kết chính, dải tần số tương ứng với liên kết chính, và tương tự.

Một cách tùy chọn, liên kết chính của thực thể ML có thể được tạo cấu hình một cách ngầm định. Cần hiểu rằng việc tạo cấu hình một cách ngầm định liên kết chính của thực thể ML giúp giảm các chi phí phát tín hiệu.

Ví dụ như, giao thức có thể quy định liên kết tương ứng với dải tần số cụ thể là liên kết chính. Ví dụ như, giao thức quy định liên kết tương ứng với dải tần số 2,4 GHz là liên kết chính.

Ví dụ như, giao thức có thể quy định: Trong nhiều liên kết được hỗ trợ bởi thực thể ML, liên kết tương ứng với dải tần số có tần số thấp nhất là liên kết chính, hoặc liên kết tương ứng với dải tần số có tần số cao nhất là liên kết chính. Ví dụ như, thực thể ML hỗ trợ dải tần số 2,4 GHz, dải tần số 5 GHz, và dải tần số 6 GHz. Khi liên kết tương ứng với dải tần số có tần số thấp nhất được sử dụng làm liên kết chính, thực thể ML sử dụng

liên kết tương ứng với dải tần số 2,4 GHz làm liên kết chính. Khi liên kết tương ứng với dải tần số có tần số cao nhất được sử dụng làm liên kết chính, thực thể ML sử dụng liên kết tương ứng với dải tần số 6 GHz làm liên kết chính.

Theo phương án này của sáng chế, trong một BSS, liên kết chính của thực thể ML AP, liên kết chính của thực thể SL, và liên kết chính của thực thể ML non-SP là giống nhau.

Có thể hiểu rằng, đối với thực thể ML, tất cả các liên kết ngoại trừ liên kết chính trong nhiều liên kết được hỗ trợ bởi thực thể ML là các liên kết không chính.

Theo sự triển khai, thực thể ML đợi chu kỳ nghỉ của kênh chính của liên kết chính để đạt tới khoảng thời gian liên khung thứ hai. Sau khi chu kỳ nghỉ của kênh chính của liên kết chính đạt tới khoảng thời gian liên khung thứ hai, mỗi lần kênh chính của liên kết chính ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1. Khi giá trị đếm của bộ đếm chờ truyền là 0, thực thể ML kết thúc quy trình chờ truyền của liên kết chính.

Cần lưu ý rằng, nếu kênh chính của liên kết chính ở trong trạng thái bận trong một khe thời gian, thực thể ML đóng băng bộ đếm chờ truyền đến khi chu kỳ nghỉ đạt tới khoảng thời gian liên khung thứ hai lần nữa.

Một cách tùy chọn, khoảng thời gian liên khung thứ hai có thể là AIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Một cách tùy chọn, kênh chính có thể là kênh chính 20 MHz. Điều này không bị giới hạn theo phương án này của sáng chế.

Một cách tùy chọn, kênh chính của liên kết chính có thể được tạo cấu hình một cách rõ ràng. Cần hiểu rằng việc tạo cấu hình một cách rõ ràng kênh chính của liên kết chính là linh hoạt.

Ví dụ như, thực thể ML có thể nhận khung MAC từ thiết bị khác, trong đó khung MAC được sử dụng để chỉ ra vị trí miền tần số của kênh chính của liên kết chính trong dải tần số tương ứng với liên kết chính. Một cách tùy chọn, khung MAC có thể là khung quản lý như là khung báo hiệu (beacon) hoặc khung đáp ứng liên hợp.

Một cách tùy chọn, kênh chính của liên kết chính có thể được tạo cấu hình một cách ngầm định. Cần hiểu rằng việc tạo cấu hình một cách ngầm định kênh chính của liên kết chính giúp giảm các chi phí phát tín hiệu.

Ví dụ như, vị trí miền tần số đặt trước của kênh chính của liên kết chính trong dải tần số tương ứng với liên kết chính có thể được quy định trong giao thức. Trong ví dụ, kênh phụ có tần số cao nhất 20 MHz trong dải tần số tương ứng với liên kết chính được sử dụng làm kênh chính. Trong ví dụ khác, kênh phụ có tần số thấp nhất 20 MHz trong dải tần số tương ứng với liên kết chính được sử dụng làm kênh chính.

S102. Khi giá trị đếm của bộ đếm chờ truyền là 0, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

K liên kết thứ nhất bao gồm liên kết chính và K-1 liên kết không chính thứ nhất, và K là số nguyên dương.

Theo phương án này của sáng chế, liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0. Cần hiểu rằng thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền của liên kết chính giảm xuống 0 là tương đương với thời điểm kết thúc của quy trình chờ truyền của liên kết chính.

Nói cách khác, đối với bất kỳ liên kết không chính, nếu liên kết không chính ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính, liên kết không chính là liên kết không chính thứ nhất. Mặt khác, liên kết không chính không là liên kết không chính thứ nhất. Tùy chọn, khoảng thời gian liên khung thứ nhất là PIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Một cách tùy chọn, trạng thái bận/ngỉ của liên kết không chính có thể được xác định dựa trên trạng thái bận/ngỉ của kênh chính của liên kết không chính. Nói cách khác, nếu kênh chính của liên kết không chính ở trong trạng thái bận, nó chỉ ra rằng liên kết không chính ở trong trạng thái bận. Nếu kênh chính của liên kết không chính ở trong trạng thái nghỉ, nó chỉ ra rằng liên kết không chính ở trong trạng thái nghỉ.

Một cách tùy chọn, kênh chính của liên kết không chính là kênh chính của 20 MHz. Điều này không bị giới hạn theo phương án này của sáng chế.

Cần hiểu rằng, khi trạng thái bận/ngỉ của liên kết không chính được xác định dựa trên trạng thái bận/ngỉ của kênh chính của liên kết không chính, kênh chính của liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0.

Một cách tùy chọn, kênh chính của liên kết không chính có thể được tạo cấu hình một cách rõ ràng. Cần hiểu rằng việc tạo cấu hình một cách rõ ràng kênh chính của liên kết không chính là linh hoạt.

Ví dụ như, thực thể ML có thể nhận khung MAC từ thiết bị khác, trong đó khung MAC được sử dụng để chỉ ra vị trí miền tần số của kênh chính của liên kết không chính trong dải tần số tương ứng với liên kết không chính. Một cách tùy chọn, khung MAC có thể là khung quản lý như là khung báo hiệu (beacon) hoặc khung đáp ứng liên hợp.

Một cách tùy chọn, kênh chính của liên kết không chính có thể được tạo cấu hình một cách ngầm định. Cần hiểu rằng việc tạo cấu hình một cách ngầm định kênh chính của liên kết không chính giúp giảm các chi phí phát tín hiệu.

Ví dụ như, vị trí miền tần số đặt trước của kênh chính của liên kết không chính trong dải tần số tương ứng với liên kết không chính có thể được quy định trong giao thức. Trong ví dụ, kênh phụ có tần số cao nhất 20 MHz trong dải tần số tương ứng với liên kết không chính được sử dụng làm kênh chính của liên kết không chính. Trong ví dụ khác, kênh phụ có tần số thấp nhất 20 MHz trong dải tần số tương ứng với liên kết không chính được sử dụng làm kênh chính của liên kết không chính.

Theo sự triển khai tùy chọn, thực thể ML gửi PPDU thứ nhất trên kênh khả dụng thứ nhất của mỗi liên kết thứ nhất trong K liên kết thứ nhất. Kênh khả dụng thứ nhất của liên kết chính bao gồm kênh chính của liên kết chính. Kênh khả dụng thứ nhất của liên kết không chính thứ nhất bao gồm kênh chính của liên kết không chính thứ nhất.

Cụ thể, đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, trước khi gửi PPDU thứ nhất, thực thể ML xác định trước tiên băng thông khả dụng của liên kết thứ nhất dựa trên trạng thái nghỉ/bận của mỗi kênh phụ trong dải tần số tương ứng với liên kết thứ nhất và yêu cầu băng thông của thực thể ML, sao cho các tài nguyên băng thông của liên kết thứ nhất được sử dụng hoàn toàn.

Cần lưu ý rằng PPDU thứ nhất là PPDU khởi tạo được gửi bởi thực thể ML trên liên kết thứ nhất. PPDU thứ nhất có thể được sử dụng để thiết lập TXOP.

Cần hiểu rằng các PPDU thứ nhất được truyền trên các liên kết thứ nhất khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ nhất khác nhau trên các liên kết thứ nhất khác nhau.

Một cách tùy chọn, PPDU thứ nhất bao gồm một trong số ba trường hợp sau.

Trường hợp 1: PPDU thứ nhất bao gồm khung MAC loại thứ nhất, nhưng không bao gồm khung MAC loại thứ hai.

Theo phương án này của sáng chế, đầu nhận không cần phản hồi lại khung đáp ứng cho khung MAC loại thứ nhất. Nói cách khác, khung MAC loại thứ nhất không cần đáp ứng.

Ví dụ như, khung MAC loại thứ nhất là khung CTS-to-self. Điều này không bị giới hạn theo phương án này của sáng chế.

Theo phương án này của sáng chế, đầu nhận cần phản hồi lại khung đáp ứng cho khung MAC loại thứ hai. Nói cách khác, khung MAC loại thứ hai cần đáp ứng.

Ví dụ như, khung MAC loại thứ hai có thể là khung RTS. Khi khung MAC loại thứ hai là khung RTS, khung đáp ứng cho khung MAC loại thứ hai là khung CTS.

Ví dụ như, khung MAC loại thứ hai có thể là khung dữ liệu. Khi khung MAC loại thứ hai là khung dữ liệu, khung đáp ứng cho khung MAC loại thứ hai là khung xác nhận (acknowledge, ACK).

Cần hiểu rằng, nếu thực thể ML gửi PPDU thứ nhất tương ứng với trường hợp 1 trên liên kết chính, thực thể ML xác nhận mặc định rằng việc thiết lập TXOP thành công.

Trường hợp 2: PPDU thứ nhất bao gồm khung MAC loại thứ hai, nhưng không bao gồm khung MAC loại thứ nhất.

Trường hợp 3: PPDU thứ nhất bao gồm khung MAC loại thứ nhất và khung MAC loại thứ hai.

Đối với trường hợp 2 hoặc trường hợp 3, giả sử rằng thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai trên liên kết chính. Nếu thực thể ML nhận khung đáp ứng cho khung MAC loại thứ hai trên liên kết chính, thực thể ML xác định rằng việc thiết lập TXOP thành công; hoặc nếu thực thể ML không nhận khung đáp ứng cho khung MAC loại thứ hai trên liên kết chính, thực thể ML xác định rằng việc thiết lập TXOP thất bại.

Với tham chiếu đến các trường hợp khác nhau của PPDU thứ nhất, phần sau mô tả cụ thể các kịch bản trong đó thực thể ML gửi PPDU thứ nhất trên K liên kết thứ nhất.

Kịch bản 1: Thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Kịch bản 2: Thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất

trên liên kết chính và một phần của các liên kết không chính thứ nhất, và gửi PPDU bao gồm khung MAC loại thứ hai trên phần khác của các liên kết không chính thứ nhất.

Trong kịch bản 1 hoặc kịch bản 2, thực thể ML xác nhận mặc định rằng việc thiết lập TXOP thành công.

Kịch bản 3: Thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Kịch bản 4: Thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai trên liên kết chính và một phần của các liên kết không chính thứ nhất, và gửi PPDU bao gồm khung MAC loại thứ nhất trên phần khác của các liên kết không chính thứ nhất.

Trong kịch bản 3 hoặc kịch bản 4, thực thể ML nhận khung đáp ứng cho khung MAC loại thứ hai trên một hoặc nhiều liên kết thứ nhất. Nếu một hoặc nhiều liên kết thứ nhất không bao gồm liên kết chính, thực thể ML xác định rằng việc thiết lập TXOP thất bại. Nếu một hoặc nhiều liên kết thứ nhất bao gồm liên kết chính, thực thể ML xác định rằng việc thiết lập TXOP thành công.

Theo phương án này của sáng chế, khi việc thiết lập TXOP thành công, thời lượng tối đa của TXOP có thể được xác định dựa trên trường thời lượng trong PPDU thứ nhất được truyền trên liên kết chính.

Theo giải pháp kỹ thuật được thể hiện trên Fig.5, vì thực thể ML đặt bộ đếm chờ truyền chỉ trên liên kết chính, khi thực hiện truy nhập kênh, thực thể ML thực hiện quy trình chờ truyền chỉ trên liên kết chính. Theo cách này, thực thể ML không thể thu được kênh thông qua sự tranh chấp trước khi quy trình chờ truyền của liên kết chính kết thúc. Điều này đảm bảo rằng xác suất thu được kênh thông qua sự tranh chấp trên liên kết chính là bằng với xác suất thu được kênh thông qua sự tranh chấp trên liên kết được hỗ trợ của thực thể SL bởi thực thể SL. Vì thế, các giải pháp kỹ thuật được đề xuất trong sáng chế có thể đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Ngoài ra, theo giải pháp kỹ thuật trên được thể hiện trên Fig.5, khi liên kết được hỗ trợ của thực thể SL và liên kết chính của thực thể ML là cùng một liên kết, thực thể SL và thực thể ML thực sự thực hiện sự tranh chấp kênh trên cùng một liên kết. Theo cách này, khi thực thể ML thu được thành công kênh thông qua sự tranh chấp trên liên kết chính, thực thể SL không gửi PPDU trên liên kết chính. Điều này đảm bảo sự đồng

bộ trong việc nhận và việc gửi được thực hiện bởi thực thể ML trên nhiều liên kết. Ví dụ như, liên kết #1 được sử dụng như liên kết chính. Khi giá trị đếm của bộ đếm chờ truyền của thực thể ML AP trên liên kết #1 là 0, thực thể ML AP gửi PPDU trên liên kết #1 và liên kết #2. Thực thể SL không gửi PPDU đến thực thể ML AP trên liên kết #1. Vì thế, thực thể ML AP có thể nhận một cách đồng bộ các tín hiệu, hoặc gửi một cách đồng bộ các tín hiệu trên liên kết #1 và liên kết #2.

Theo phương án tùy chọn, dựa trên phương pháp truyền thông được thể hiện trên Fig.5, trên Fig.6, khi thực thể ML thiết lập một cách thành công TXOP, phương pháp truyền thông còn bao gồm các bước S103 và S104.

S103. Thực thể ML xác định N liên kết thứ hai tương ứng với TXOP từ K liên kết thứ nhất.

N liên kết thứ hai bao gồm liên kết chính và N-1 liên kết không chính thứ hai, và N là số nguyên dương nhỏ hơn hoặc bằng K.

Theo phương án này của sáng chế, liên kết không chính thứ hai là liên kết không chính thứ nhất mà đáp ứng điều kiện đặt trước.

Một cách tùy chọn, điều kiện đặt trước bao gồm một trong số phần sau.

Điều kiện 1: Trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất.

Điều kiện 2: Trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai, và nhận khung đáp ứng cho khung MAC loại thứ hai.

S104. Thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ hai trong N liên kết thứ hai.

PPDU thứ hai là khác với PPDU thứ nhất. Nói cách khác, PPDU thứ hai là PPDU khác với PPDU thứ nhất.

Cần hiểu rằng các PPDU thứ hai được gửi bởi thực thể ML trên các liên kết thứ hai khác nhau có thể là các PPDU khác nhau, để triển khai thông lượng cực cao.

Theo phương án này của sáng chế, thực thể ML gửi PPDU thứ hai trên một liên kết thứ hai. Nếu thực thể ML không nhận khung đáp ứng trên liên kết thứ hai ở trong thời lượng cụ thể, nó chỉ ra rằng sự truyền PPDU thứ hai trên liên kết thứ hai thất bại. Ví dụ như, khung đáp ứng có thể là khung BA. Điều này không bị giới hạn theo phương án này của sáng chế.

Cần hiểu rằng, trong kịch bản trong đó sự truyền PPDU thứ hai thất bại trên liên kết thứ hai, nếu thực thể ML không thực hiện việc xử lý tương ứng trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thất bại, mà tiếp tục gửi PPDU thứ hai thất bại, PPDU thứ hai được gửi bởi thực thể ML có thể luôn thất bại truyền, mà ảnh hưởng đến truyền thông thích hợp của thực thể ML.

Phần sau mô tả cách xử lý được sử dụng bởi thực thể ML trong kịch bản trong đó sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai.

Cách xử lý 1: Nếu sự truyền PPDU thứ hai thành công trên liên kết chính và sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết không chính thứ hai, thực thể ML dừng gửi PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thất bại, và tiếp tục gửi, đến khi TXOP kết thúc, và PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thành công.

Ví dụ như, thực thể ML gửi PPDU thứ hai trên liên kết không chính #1, liên kết không chính #2, liên kết không chính #3, và liên kết chính riêng biệt. Nếu sự truyền PPDU thứ hai thất bại trên liên kết không chính #2, thực thể ML dừng gửi PPDU thứ hai trên liên kết không chính #2, và tiếp tục gửi PPDU thứ hai trên liên kết không chính #1, liên kết không chính #3, và liên kết chính.

Cách xử lý 2: Nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML dừng gửi PPDU thứ hai trên N liên kết thứ hai. Sau đó, thực thể ML đợi chu kỳ nghỉ của liên kết chính để đạt tới khoảng thời gian liên khung thứ nhất. Khi chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba.

P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, và P là số nguyên dương nhỏ hơn hoặc bằng N. Liên kết không chính thứ ba là liên kết không chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất, và thời điểm thứ nhất là thời điểm mà tại đó chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất.

Ví dụ được đề xuất cho sự mô tả với tham chiếu đến Fig.7, thực thể ML gửi PPDU thứ hai #1 trên liên kết không chính #1, liên kết không chính #2, và liên kết chính riêng biệt. Vì thực thể ML không nhận khung BA trên liên kết không chính #1, thực thể ML xác định rằng sự truyền PPDU thứ hai #1 thất bại trên liên kết không chính #1. Trong

trường hợp này, thực thể ML tạm dừng gửi PPDU thứ hai trên liên kết không chính #1, liên kết không chính #2, và liên kết chính. Sau một PIFS, liên kết không chính #1 và liên kết chính trong PIFS ở trong trạng thái nghỉ, và liên kết không chính #2 ở trong trạng thái bận. Vì thế, thực thể ML có thể xác định rằng liên kết không chính #1 và liên kết chính là các liên kết thứ ba. Trong trường hợp này, thực thể ML gửi PPDU thứ hai #2 trên liên kết không chính #1 và liên kết chính, và thực thể ML không gửi PPDU thứ hai #2 trên liên kết không chính #2.

Cách xử lý 3: Nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML dừng gửi PPDU thứ hai trên N liên kết thứ hai. Sau đó, thực thể ML thực hiện quy trình chờ truyền trên liên kết chính. Khi quy trình chờ truyền của liên kết chính kết thúc, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba.

P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, và P là số nguyên dương nhỏ hơn hoặc bằng N. Liên kết không chính thứ ba là liên kết không chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính.

Cần hiểu rằng đối với các chi tiết về việc thực thể ML thực hiện quy trình chờ truyền trên liên kết chính, tham chiếu đến sự mô tả trên theo bước S101. Các chi tiết không được mô tả lần nữa ở đây.

Ví dụ được đề xuất cho sự mô tả với tham chiếu đến Fig.8, thực thể ML gửi PPDU thứ hai #1 trên liên kết không chính #1, liên kết không chính #2, và liên kết chính riêng biệt. Vì thực thể ML không nhận khung BA trên liên kết không chính #1, thực thể ML xác định rằng sự truyền PPDU thứ hai #1 thất bại trên liên kết không chính #1. Trong trường hợp này, thực thể ML tạm dừng gửi PPDU thứ hai trên liên kết không chính #1, liên kết không chính #2, và liên kết chính. Thực thể ML đặt giá trị đếm của bộ đếm chờ truyền của liên kết chính. Trong PIFS trước thời điểm mà tại đó bộ đếm chờ truyền của liên kết chính giảm xuống 0, liên kết không chính #1 và liên kết chính ở trong trạng thái nghỉ, và liên kết không chính #2 ở trong trạng thái bận. Vì thế, thực thể ML có thể xác định rằng liên kết không chính #1 và liên kết chính là các liên kết thứ ba. Trong trường hợp này, thực thể ML gửi PPDU thứ hai #2 trên liên kết không chính #1 và liên kết chính, và thực thể ML không gửi PPDU thứ hai #2 trên liên kết không chính #2.

Cần hiểu rằng, theo cách xử lý 2 hoặc cách xử lý 3 trên, đối với mỗi liên kết thứ ba trong P liên kết thứ ba, việc thực thể ML gửi PPDU thứ hai trên các liên kết thứ ba bao gồm: Thực thể ML gửi PPDU thứ hai trên các kênh khả dụng thứ hai của các liên kết thứ ba. Đối với một liên kết, kênh khả dụng thứ hai là tập con của các kênh khả dụng thứ nhất. Các kênh khả dụng thứ hai cũng bao gồm kênh chính.

Cần hiểu rằng, theo cách xử lý 2 hoặc cách xử lý 3 trên, việc sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai có nghĩa cụ thể rằng sự truyền PPDU thứ hai thất bại trên liên kết chính, và/hoặc sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết không chính thứ hai.

Theo bất kỳ một trong các cách xử lý trên, trong kịch bản trong đó sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, truyền thông thích hợp của thực thể ML có thể được đảm bảo.

Fig.9(a) thể hiện phương pháp truyền thông theo phương án của sáng chế. Phương pháp bao gồm các bước sau.

S201. Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Thực thể ML hỗ trợ K liên kết thứ nhất, và K là số nguyên dương lớn hơn hoặc bằng 2. Bộ đếm chờ truyền được bố trí trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, quy trình chờ truyền của liên kết thứ nhất bao gồm bước sau: Thực thể ML đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung thứ hai. Sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1. Khi giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất là 0, thực thể ML kết thúc quy trình chờ truyền của liên kết thứ nhất.

Theo phương án này của sáng chế, nếu liên kết thứ nhất ở trong trạng thái bận trong một khe thời gian, thực thể ML đóng băng bộ đếm chờ truyền của liên kết thứ nhất đến khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai lần nữa. Cần hiểu rằng, việc bộ đếm chờ truyền của liên kết thứ nhất bị đóng băng tương đương với việc quy trình chờ truyền của liên kết thứ nhất bị tạm dừng.

Khoảng thời gian liên khung thứ hai có thể là AIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Trạng thái bận/ngỉ của liên kết thứ nhất có thể được xác định dựa trên trạng thái bận/ngỉ của kênh chính của liên kết thứ nhất. Nói cách khác, nếu kênh chính của liên kết thứ nhất ở trong trạng thái bận, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái bận. Nếu kênh chính của liên kết thứ nhất ở trong trạng thái nghỉ, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái nghỉ.

Một cách tùy chọn, kênh chính của liên kết thứ nhất có thể là kênh chính 20 MHz.

Một cách tùy chọn, kênh chính của liên kết thứ nhất có thể được tạo cấu hình một cách rõ ràng. Cần hiểu rằng việc tạo cấu hình một cách rõ ràng kênh chính của liên kết thứ nhất là linh hoạt.

Ví dụ như, thực thể ML có thể nhận khung MAC từ thiết bị khác, trong đó khung MAC được sử dụng để chỉ ra vị trí miền tần số của kênh chính của liên kết thứ nhất trong dải tần số tương ứng với liên kết thứ nhất. Một cách tùy chọn, khung MAC có thể là khung quản lý như là khung báo hiệu hoặc khung đáp ứng liên hợp.

Một cách tùy chọn, kênh chính của liên kết thứ nhất có thể được tạo cấu hình một cách ngầm định. Cần hiểu rằng việc tạo cấu hình một cách ngầm định kênh chính của liên kết thứ nhất giúp giảm các chi phí phát tín hiệu.

Ví dụ như, vị trí miền tần số đặt trước của kênh chính của liên kết thứ nhất trong dải tần số tương ứng với liên kết thứ nhất có thể được quy định trong giao thức. Trong ví dụ, kênh phụ có tần số cao nhất 20 MHz trong dải tần số tương ứng với liên kết thứ nhất được sử dụng làm kênh chính của liên kết thứ nhất. Trong ví dụ khác, kênh phụ có tần số thấp nhất 20 MHz trong dải tần số tương ứng với liên kết thứ nhất được sử dụng làm kênh chính của liên kết thứ nhất.

Theo phương án này của sáng chế, khi thực thể ML thực hiện riêng biệt các quy trình chờ truyền của K liên kết thứ nhất, nếu quy trình chờ truyền của một liên kết kết thúc trước tiên, thực thể ML thực hiện bước S202.

S202. Khi quy trình chờ truyền của liên kết mục tiêu kết thúc, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Liên kết mục tiêu là liên kết thứ nhất mà quy trình chờ truyền kết thúc trước tiên trong K liên kết thứ nhất. Nói cách khác, liên kết mục tiêu là liên kết thứ nhất mà ở trong

K liên kết thứ nhất và trên đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0 trước tiên.

N liên kết thứ hai bao gồm liên kết mục tiêu và N-1 liên kết khả dụng, và N là số nguyên dương nhỏ hơn hoặc bằng K.

Theo phương án này của sáng chế, liên kết khả dụng là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu. Nói cách khác, liên kết khả dụng là liên kết thứ nhất mà kênh chính ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu.

Cần hiểu rằng, nếu một liên kết thứ nhất (hoặc kênh chính của liên kết thứ nhất) ở trong trạng thái bận trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, liên kết thứ nhất không là liên kết khả dụng.

Theo phương án này của sáng chế, sau khi quy trình chờ truyền của liên kết mục tiêu kết thúc, thực thể ML dừng quy trình chờ truyền của liên kết thứ nhất khác khác với liên kết mục tiêu trong K liên kết thứ nhất đến khi TXOP kết thúc.

Cần hiểu rằng các PPDU thứ nhất được truyền trên các liên kết thứ hai khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ nhất khác nhau trên các liên kết thứ hai khác nhau.

S203. Nếu sự truyền PPDU thứ nhất thất bại trong một hoặc nhiều liên kết thứ hai, thực thể ML bỏ qua việc gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ nhất thất bại.

PPDU thứ hai và PPDU thứ nhất là hai PPDU khác nhau. Nói cách khác, PPDU thứ hai là PPDU khác với PPDU thứ nhất.

Ví dụ như, việc truyền của PPDU thứ nhất thất bại trên liên kết thứ hai có thể tham chiếu đến việc thực thể ML không nhận, trên liên kết thứ hai, khung đáp ứng tương ứng với PPDU thứ nhất. Cần hiểu rằng khung đáp ứng tương ứng với PPDU thứ nhất được sử dụng để đáp ứng cho khung MAC được mang trong PPDU thứ nhất. Ví dụ như, nếu PPDU thứ nhất mang khung RTS, khung đáp ứng cho PPDU thứ nhất có thể là khung CTS.

Một cách tùy chọn, chu kỳ thời gian đặt trước có thể được tạo cấu hình trước hoặc

được quy định trong giao thức. Điều này không bị giới hạn theo phương án này của sáng chế.

Ví dụ như, thực thể ML gửi riêng biệt PPDU thứ nhất trên liên kết #1, liên kết #3, và liên kết #4. Nếu thực thể ML không nhận trên liên kết #1 khung đáp ứng tương ứng với PPDU thứ nhất, thực thể ML có thể xác định rằng sự truyền PPDU thứ nhất thất bại trên liên kết #1. Vì thế, thực thể ML không gửi PPDU thứ hai trên liên kết #1 ở trong chu kỳ thời gian đặt trước.

Một cách tùy chọn, trên Fig.9(b), bước S203 trên Fig.9(a) có thể được thay thế bằng bước S204.

S204. Nếu sự truyền PPDU thứ nhất thất bại trên một hoặc nhiều liên kết thứ hai, thực thể ML bỏ qua việc gửi PPDU thứ hai trên liên kết thứ hai ở trong chu kỳ thời gian đặt trước.

Ví dụ như, thực thể ML gửi riêng biệt PPDU thứ nhất trên liên kết #1, liên kết #3, và liên kết #4. Nếu thực thể ML không nhận trên liên kết #1 khung đáp ứng tương ứng với PPDU thứ nhất, thực thể ML có thể xác định rằng sự truyền PPDU thứ nhất thất bại trên liên kết #1. Vì thế, thực thể ML không gửi PPDU thứ hai trên liên kết #1, liên kết #3, và liên kết #4 ở trong chu kỳ thời gian đặt trước.

Theo giải pháp kỹ thuật được thể hiện trên Fig.9(a) hoặc Fig.9(b), khi thực thể ML truyền thất bại PPDU thứ nhất trên một hoặc nhiều liên kết thứ hai, thực thể ML bị cấm gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thất bại, hoặc thực thể ML bị cấm gửi PPDU thứ hai trên N liên kết thứ hai ở trong chu kỳ thời gian đặt trước. Theo cách này, ở trong chu kỳ thời gian đặt trước, thực thể ML không thể sử dụng nhiều liên kết (ví dụ như, N liên kết thứ hai hoặc các liên kết thứ hai mà trên đó sự truyền PPDU thứ nhất thất bại). Nếu một liên kết mà thực thể ML không thể sử dụng và ở trong nhiều liên kết được hỗ trợ bởi thực thể SL, ở trong chu kỳ thời gian đặt trước, vì thực thể ML không thể thực hiện sự tranh chấp kênh trên liên kết được hỗ trợ bởi thực thể SL, xác suất mà thực thể SL thu được kênh thông qua sự tranh chấp gia tăng. Điều này đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Fig.10 thể hiện phương pháp truyền thông theo phương án của sáng chế. Phương pháp bao gồm các bước sau.

S301. Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Thực thể ML hỗ trợ K liên kết thứ nhất, và K là số nguyên lớn hơn hoặc bằng 2. Bộ đếm chờ truyền được bố trí trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, quy trình chờ truyền của liên kết thứ nhất bao gồm bước sau: Thực thể ML đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung thứ hai. Sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1. Khi giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất là 0, thực thể ML kết thúc quy trình chờ truyền của liên kết thứ nhất.

Theo phương án này của sáng chế, nếu liên kết thứ nhất ở trong trạng thái bận trong một khe thời gian, thực thể ML đóng băng bộ đếm chờ truyền của liên kết thứ nhất đến khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai lần nữa. Cần hiểu rằng, việc bộ đếm chờ truyền của liên kết thứ nhất bị đóng băng tương đương với việc quy trình chờ truyền của liên kết thứ nhất bị tạm dừng.

Khoảng thời gian liên khung thứ hai có thể là AIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Trạng thái bận/ngỉ của liên kết thứ nhất có thể được xác định dựa trên trạng thái bận/ngỉ của kênh chính của liên kết thứ nhất. Nói cách khác, nếu kênh chính của liên kết thứ nhất ở trong trạng thái bận, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái bận. Nếu kênh chính của liên kết thứ nhất ở trong trạng thái nghỉ, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái nghỉ.

Một cách tùy chọn, kênh chính có thể là kênh chính 20 MHz. Đối với mỗi liên kết thứ nhất, đối với phương pháp để tạo cấu hình kênh chính của liên kết thứ nhất, tham chiếu đến sự mô tả trên. Các chi tiết không được mô tả lần nữa ở đây.

S302. Thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai.

N liên kết thứ hai là tập con của K liên kết thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng K.

Theo phương án này của sáng chế, liên kết thứ hai là liên kết thứ nhất mà quy

trình chờ truyền đã kết thúc và ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất.

Nói cách khác, nếu quy trình chờ truyền của một liên kết thứ nhất chưa kết thúc trước thời điểm thứ nhất, liên kết thứ nhất không là liên kết thứ hai. Ngoài ra, nếu liên kết thứ nhất ở trong trạng thái bận trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất, liên kết thứ nhất không là liên kết thứ hai.

Ví dụ như, thực thể ML thực hiện quy trình chờ truyền trên liên kết #1, liên kết #2, liên kết #3, và liên kết #4 riêng biệt. Trước thời điểm thứ nhất, quy trình chờ truyền của liên kết #1 đã kết thúc, quy trình chờ truyền của liên kết #2 đã kết thúc, và quy trình chờ truyền của liên kết #4 đã kết thúc. Ngoài ra, trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất, liên kết #1 ở trong trạng thái nghỉ, liên kết #2 ở trong trạng thái bận, và liên kết #4 ở trong trạng thái nghỉ. Vì thế, thực thể ML có thể xác định rằng liên kết #1 và liên kết #4 là các liên kết thứ hai.

Một cách tùy chọn, thời điểm thứ nhất có thể được tạo cấu hình trước hoặc được quy định trong giao thức.

Tùy chọn, thời điểm thứ nhất có thể là thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu.

Ví dụ như, liên kết mục tiêu có thể là liên kết thứ hai mà quy trình chờ truyền kết thúc cuối cùng trong N liên kết thứ hai.

Ví dụ như, liên kết mục tiêu có thể là liên kết thứ nhất thứ k trong K liên kết thứ nhất để kết thúc quy trình chờ truyền, và k là số nguyên lớn hơn 1 và nhỏ hơn hoặc bằng K.

Cần hiểu rằng các PPDU thứ nhất được truyền trên các liên kết thứ hai khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ nhất khác nhau trên các liên kết thứ hai khác nhau.

Theo giải pháp kỹ thuật trên được thể hiện trên Fig.10, mặc dù ML thực hiện quy trình chờ truyền trên tất cả K liên kết thứ nhất, liên kết thứ hai được sử dụng để gửi PPDU thứ nhất cần thỏa mãn điều kiện mà quy trình chờ truyền của liên kết thứ hai đã kết thúc. Nói cách khác, trên một liên kết, thực thể ML có thể thu được kênh thông qua sự tranh chấp chỉ sau khi thực thể ML kết thúc quy trình chờ truyền trên liên kết. So sánh với công nghệ thông thường trong đó thực thể ML có thể thu được kênh thông qua

sự tranh chấp trên một liên kết kể cả nếu quy trình chờ truyền trên liên kết chưa kết thúc, giải pháp kỹ thuật trong sáng chế giảm xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp trên một liên kết. Điều này đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Fig.11(a) thể hiện phương pháp truyền thông theo phương án của sáng chế. Phương pháp bao gồm các bước sau.

S401. Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Thực thể ML hỗ trợ K liên kết thứ nhất, và K là số nguyên lớn hơn hoặc bằng 2. Bộ đếm chờ truyền được bố trí trên mỗi liên kết thứ nhất trong K liên kết thứ nhất.

Đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, quy trình chờ truyền trên liên kết thứ nhất bao gồm bước sau: Thực thể ML đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung thứ hai. Sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1.

Khoảng thời gian liên khung thứ hai có thể là AIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Theo phương án này của sáng chế, phạm vi giá trị của bộ đếm chờ truyền của liên kết thứ nhất bao gồm các số nguyên âm. Nói cách khác, sau khi thực thể ML giảm giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất xuống 0, thực thể ML không kết thúc quy trình chờ truyền của liên kết thứ nhất, mà tiếp tục chờ truyền.

Ví dụ như, giá trị khởi tạo của bộ đếm chờ truyền của liên kết #1 là 5. Sau khi chu kỳ nghỉ của liên kết #1 đạt tới khoảng thời gian liên khung thứ hai, nếu liên kết #1 ở trong trạng thái nghỉ trong sáu khe thời gian liên tiếp, giá trị đếm của bộ đếm chờ truyền của liên kết #1 có thể là -1.

Theo phương án này của sáng chế, nếu liên kết thứ nhất ở trong trạng thái bận trong một khe thời gian, thực thể ML đóng băng bộ đếm chờ truyền của liên kết thứ nhất đến khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai lần nữa. Cần hiểu rằng, việc bộ đếm chờ truyền của liên kết thứ nhất bị đóng băng tương đương với việc quy trình chờ truyền của liên kết thứ nhất bị tạm dừng.

Trạng thái bận/ngỉ của liên kết thứ nhất có thể được xác định dựa trên trạng thái

bận/ngủ của kênh chính của liên kết thứ nhất. Nói cách khác, nếu kênh chính của liên kết thứ nhất ở trong trạng thái bận, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái bận. Nếu kênh chính của liên kết thứ nhất ở trong trạng thái nghỉ, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái nghỉ.

Một cách tùy chọn, kênh chính có thể là kênh chính 20 MHz. Đối với mỗi liên kết thứ nhất, đối với phương pháp để tạo cấu hình kênh chính của liên kết thứ nhất, tham chiếu đến sự mô tả trên. Các chi tiết không được mô tả lần nữa ở đây.

S402. Khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai.

N liên kết thứ hai là tập con của K liên kết thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng K.

Theo phương án này của sáng chế, liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ hai trước thời điểm hiện tại. Nói cách khác, liên kết thứ hai là liên kết thứ nhất mà trên đó bộ đếm chờ truyền không bị đóng băng. Nói cách khác, liên kết thứ hai là liên kết thứ nhất mà trên đó quy trình chờ truyền không bị tạm dừng.

Theo bước S402, thời điểm hiện tại là thời điểm mà tại đó tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0.

Tùy chọn, trong một khe thời gian, thực thể ML có thể thu thập các số liệu thống kê về tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất, để xác định tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất có nhỏ hơn hoặc bằng 0 hay không.

Tùy chọn, thực thể ML còn được tạo cấu hình với bộ đếm mục tiêu, và bộ đếm mục tiêu được tạo cấu hình để ghi lại tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất. Theo cách này, trong một khe thời gian, ML có thể xác định, bằng cách xác định tổng các giá trị đếm của các bộ đếm mục tiêu có nhỏ hơn hoặc bằng 0 hay không, tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất có nhỏ hơn hoặc bằng 0 hay không.

Trong khi triển khai cụ thể, thực thể ML tạo cấu hình bộ đếm mục tiêu, và giá trị khởi tạo của bộ đếm mục tiêu bằng với tổng các giá trị khởi tạo của các bộ đếm chờ

truyền của K liên kết thứ nhất. Đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần kênh chính của liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm mục tiêu đi 1. Nói cách khác, mỗi lần thực thể ML giảm bộ đếm chờ truyền của một liên kết thứ nhất đi 1, thực thể ML giảm giá trị đếm của bộ đếm mục tiêu đi 1.

Một cách tùy chọn, trên Fig.11(b), bước S402 trên Fig.11(a) có thể được thay thế bằng bước S403.

S403. Khi tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Theo phương án này của sáng chế, liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ hai trước thời điểm hiện tại. Nói cách khác, liên kết thứ hai là liên kết thứ nhất mà trên đó bộ đếm chờ truyền không bị đóng băng. Nói cách khác, liên kết thứ hai là liên kết thứ nhất mà trên đó quy trình chờ truyền không bị tạm dừng.

Theo bước S403, thời điểm hiện tại là thời điểm mà tại đó tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ hai nhỏ hơn hoặc bằng 0.

Tùy chọn, trong một khe thời gian, thực thể ML có thể thu thập các số liệu thống kê về tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai, để xác định tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ nhất có nhỏ hơn hoặc bằng 0 hay không.

Cần hiểu rằng, theo bước S403 hoặc bước S402, các PPDU thứ nhất được truyền trên các liên kết thứ hai khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ nhất khác nhau trên các liên kết thứ hai khác nhau.

Theo giải pháp kỹ thuật được thể hiện trên Fig.11(a) hoặc Fig.11(b), mặc dù thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, thực thể ML có thể thu được một cách thành công kênh thông qua sự tranh chấp chỉ khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0 hoặc tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0. Nói cách khác, đối với thực thể ML, các giá trị đếm của các bộ

đếm chờ truyền của một hoặc nhiều liên kết thứ nhất cần phải nhỏ hơn 0. Điều này yêu cầu một hoặc nhiều liên kết thứ nhất nghỉ trong thời gian tương đối dài. Theo cách này, xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp bị giảm. Việc xác suất mà thực thể ML thu được kênh thông qua sự tranh chấp bị giảm làm suy yếu lợi thế của thực thể ML so với thực thể SL trong sự tranh chấp kênh, đảm bảo công bằng cho thực thể SL trong sự tranh chấp kênh, và vì thế đảm bảo truyền thông thích hợp của thực thể SL.

Fig.12 thể hiện phương pháp truyền thông theo phương án của sáng chế. Phương pháp bao gồm các bước sau.

S501. Thực thể ML thực hiện quy trình chờ truyền trên liên kết thứ nhất.

Thực thể ML hỗ trợ nhiều liên kết.

Theo phương án này của sáng chế, bộ đếm chờ truyền có thể được bố trí trên mỗi liên kết trong nhiều liên kết. Tuy nhiên, mỗi lần thực thể ML khởi tạo truy nhập kênh của nhiều liên kết, thực thể ML thực hiện chờ truyền bằng cách sử dụng chỉ bộ đếm chờ truyền của một liên kết (cụ thể, liên kết thứ nhất).

Một cách tùy chọn, một liên kết ngẫu nhiên trong nhiều liên kết đóng vai trò là liên kết thứ nhất. Ngoài ra, nhiều liên kết mỗi liên kết đóng vai trò là liên kết thứ nhất đến lượt theo thứ tự tuần hoàn đặt trước.

Một cách tùy chọn, thứ tự tuần hoàn có thể là thứ tự giảm dần của các số trình tự của nhiều liên kết, thứ tự lớn dần của các số trình tự của nhiều liên kết, hoặc thứ tự giả ngẫu nhiên của các số trình tự của nhiều liên kết. Điều này không bị giới hạn theo phương án này của sáng chế.

Ví dụ như, thực thể ML hỗ trợ bảy liên kết, và số trình tự của bảy liên kết là 0, 1, 2, 3, 4, 5, và 6. Thứ tự tuần hoàn đặt trước là “01204465”, và mỗi chữ số trong thứ tự tuần hoàn là số trình tự của liên kết. Theo cách này, khi thực hiện truy nhập kênh lần thứ nhất, thực thể ML sử dụng liên kết mà số trình tự là 0 làm liên kết thứ nhất. Khi thực hiện truy nhập kênh lần thứ hai, thực thể ML sử dụng liên kết mà số trình tự là 1 làm liên kết thứ nhất. Bằng cách tương tự, khi thực hiện truy nhập kênh lần thứ mười, thực thể ML sử dụng liên kết mà số trình tự là 1 làm liên kết thứ nhất.

Cần hiểu rằng các thực thể ML khác nhau có thể được tạo cấu hình với các thứ tự tuần hoàn khác nhau. Điều này không bị hạn chế trong sáng chế.

Theo phương án này của sáng chế, việc thực thể ML thực hiện quy trình chờ truyền trên liên kết thứ nhất bao gồm bước sau: Thực thể ML đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung thứ hai. Sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, thực thể ML giảm giá trị đếm của bộ đếm chờ truyền đi 1. Khi giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất là 0, thực thể ML kết thúc quy trình chờ truyền của liên kết thứ nhất.

Theo phương án này của sáng chế, nếu liên kết thứ nhất ở trong trạng thái bận trong một khe thời gian, thực thể ML đóng băng bộ đếm chờ truyền của liên kết thứ nhất đến khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai lần nữa. Cần hiểu rằng, việc bộ đếm chờ truyền của liên kết thứ nhất bị đóng băng tương đương với việc quy trình chờ truyền của liên kết thứ nhất bị tạm dừng.

Khoảng thời gian liên khung thứ hai có thể là AIFS. Điều này không bị giới hạn theo phương án này của sáng chế.

Trạng thái bận/ngỉ của liên kết thứ nhất có thể được xác định dựa trên trạng thái bận/ngỉ của kênh chính của liên kết thứ nhất. Nói cách khác, nếu kênh chính của liên kết thứ nhất ở trong trạng thái bận, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái bận. Nếu kênh chính của liên kết thứ nhất ở trong trạng thái nghỉ, nó chỉ ra rằng liên kết thứ nhất ở trong trạng thái nghỉ.

Một cách tùy chọn, kênh chính của liên kết thứ nhất có thể là kênh chính 20 MHz. Đối với phương pháp để tạo cấu hình kênh chính của liên kết thứ nhất, tham chiếu đến sự mô tả trên. Các chi tiết không được mô tả lần nữa ở đây.

S502. Khi quy trình chờ truyền của liên kết thứ nhất kết thúc, thực thể ML gửi PPDU thứ nhất trên mỗi liên kết khả dụng trong N liên kết khả dụng.

N liên kết khả dụng bao gồm liên kết thứ nhất và N-1 liên kết thứ hai, và N là số nguyên dương.

Cần hiểu rằng liên kết thứ hai và liên kết thứ nhất là hai liên kết khác nhau. Liên kết thứ hai ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết thứ nhất.

Cần hiểu rằng các PPDU thứ nhất được truyền trên các liên kết thứ hai khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ nhất khác

nhau trên các liên kết thứ hai khác nhau.

Theo giải pháp kỹ thuật được thể hiện trên Fig.12, mỗi lần truy nhập kênh được thực hiện, thực thể ML thực hiện quy trình chờ truyền chỉ trên liên kết thứ nhất. Nói cách khác, thực thể ML thực hiện sự tranh chấp kênh trên chỉ một liên kết duy nhất. Xác suất thực thể ML thu được kênh thông qua sự tranh chấp trên một liên kết bằng với xác suất thực thể SL thu được kênh thông qua sự tranh chấp trên một liên kết. Theo cách này, công bằng cho thực thể SL trong sự tranh chấp kênh được đảm bảo, và vì thế truyền thông thích hợp của thực thể SL được đảm bảo.

Theo các giải pháp kỹ thuật được thể hiện trên Fig.9(a), Fig.9(b), Fig.10, Fig.11(a), Fig.11(b), hoặc Fig.12, PPDU có thể bao gồm ba trường hợp sau.

Trường hợp 1: PPDU thứ nhất bao gồm khung MAC loại thứ nhất, nhưng không bao gồm khung MAC loại thứ hai.

Trường hợp 2: PPDU thứ nhất bao gồm khung MAC loại thứ hai, nhưng không bao gồm khung MAC loại thứ nhất.

Trường hợp 3: PPDU thứ nhất bao gồm khung MAC loại thứ nhất và khung MAC loại thứ hai.

Đối với các mô tả chi tiết của khung MAC loại thứ nhất và khung MAC loại thứ hai, tham chiếu đến các mô tả theo bước S102. Các chi tiết không được mô tả lần nữa ở đây.

Với tham chiếu đến các trường hợp khác nhau của PPDU thứ nhất, phần sau mô tả cụ thể các kịch bản trong đó thực thể ML gửi PPDU thứ nhất trên K liên kết thứ hai.

Kịch bản 1: Thực thể ML gửi PPDU thứ nhất bao gồm chỉ khung MAC loại thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Kịch bản 2: Thực thể ML gửi PPDU thứ nhất bao gồm chỉ khung MAC loại thứ nhất trên một phần của liên kết thứ hai, và gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai trên phần khác của liên kết thứ hai.

Trong kịch bản 1 hoặc kịch bản 2, thực thể ML xác nhận mặc định rằng việc thiết lập TXOP thành công.

Kịch bản 3: Thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Trong kịch bản 3, nếu thực thể ML không nhận khung đáp ứng cho khung MAC

loại thứ hai trên bất kỳ liên kết thứ hai, thực thể ML xác nhận rằng việc thiết lập TXOP thất bại. Nếu thực thể ML nhận khung đáp ứng cho khung MAC loại thứ hai trên ít nhất một liên kết thứ hai, thực thể ML xác nhận rằng việc thiết lập TXOP thành công.

Theo phương án tùy chọn, dựa trên các giải pháp kỹ thuật được thể hiện trên Fig.10 đến Fig.12, trên Fig.13, khi thực thể ML thiết lập một cách thành công TXOP, phương pháp truyền thông còn bao gồm các bước S601 và S602 sau.

S601. Thực thể ML xác định P liên kết thứ ba tương ứng với TXOP từ N liên kết thứ hai.

P liên kết thứ ba là tập con của N liên kết thứ hai, và P là số nguyên dương nhỏ hơn hoặc bằng N.

Theo phương án này của sáng chế, liên kết thứ ba là liên kết thứ hai mà đáp ứng điều kiện đặt trước.

Một cách tùy chọn, điều kiện đặt trước bao gồm một trong số phần sau.

Điều kiện 1: Trên liên kết thứ hai, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất.

Điều kiện 2: Trên liên kết thứ hai, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai, và nhận khung đáp ứng cho khung MAC loại thứ hai.

S602. Thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba.

Cần hiểu rằng các PPDU thứ hai được truyền trên các liên kết thứ ba khác nhau có thể là khác nhau. Nói cách khác, thực thể ML có thể gửi các PPDU thứ hai khác nhau trên các liên kết thứ ba khác nhau.

Một cách tùy chọn, khi sự truyền PPDU thất bại trên một hoặc nhiều liên kết thứ ba, thực thể ML có thể sử dụng bất kỳ một trong số các cách xử lý sau.

Cách xử lý 1: Thực thể ML dừng gửi PPDU thứ hai thất bại trên liên kết thứ ba mà trên đó sự truyền PPDU thứ hai thất bại, và tiếp tục gửi, đến khi TXOP kết thúc, và PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thành công.

Cách xử lý 2: Thực thể ML dừng gửi PPDU thứ hai trên P liên kết thứ ba. Thực thể ML đợi đến thời điểm đặt trước để xác định L liên kết thứ tư. Ngoài ra, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ tư trong L liên kết thứ tư.

L liên kết thứ tư là tập hợp con của P liên kết thứ ba, và L là số nguyên dương

nhỏ hơn hoặc bằng P. Liên kết thứ tư là liên kết thứ ba mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm đặt trước.

Cần hiểu rằng thời điểm đặt trước có thể được tạo cấu hình trước hoặc được quy định trong giao thức.

Cách xử lý 3: Thực thể ML dùng gửi PPDU thứ hai trên P liên kết thứ ba. Thực thể ML thực hiện quy trình chờ truyền trên mỗi liên kết thứ ba trong P liên kết thứ ba. Khi quy trình chờ truyền của liên kết thứ ba mục tiêu kết thúc, thực thể ML xác định L liên kết thứ tư. Ngoài ra, thực thể ML gửi PPDU thứ hai trên mỗi liên kết thứ tư trong L liên kết thứ tư.

L liên kết thứ tư là tập hợp con của P liên kết thứ ba, và L là số nguyên dương nhỏ hơn hoặc bằng P. Liên kết thứ tư là liên kết thứ ba mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước khi quy trình chờ truyền của liên kết thứ ba mục tiêu kết thúc. Liên kết thứ ba mục tiêu có thể là liên kết thứ ba mà quy trình chờ truyền kết thúc trước tiên trong P liên kết thứ ba.

Theo bất kỳ một trong các cách xử lý trên, trong kịch bản trong đó sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ ba, truyền thông thích hợp của thực thể ML có thể được đảm bảo.

Phần trên chủ yếu mô tả các giải pháp được đề xuất theo các phương án của sáng chế từ góc độ của thực thể ML. Có thể hiểu rằng, để triển khai các chức năng trên, thực thể ML bao gồm cấu trúc phần cứng tương ứng và/hoặc môđun phần mềm cho việc triển khai mỗi chức năng. Người có hiểu biết trung bình trong lĩnh vực có thể dễ dàng nhận thức rằng, trong sự kết hợp với các đơn vị và các bước thuật toán của các ví dụ được mô tả theo các phương án được bộc lộ trong bản mô tả, sáng chế có thể được triển khai bằng phần cứng hoặc sự kết hợp của phần cứng và phần mềm máy tính. Việc chức năng được thực hiện bởi phần cứng hoặc phần cứng được chạy bởi phần mềm máy tính phụ thuộc vào các ứng dụng cụ thể và các điều kiện ràng buộc thiết kế của các giải pháp kỹ thuật. Người có hiểu biết trung bình trong lĩnh vực có thể sử dụng phương pháp khác để triển khai các chức năng được mô tả cho mỗi ứng dụng cụ thể, nhưng không được coi rằng sự triển khai vượt quá phạm vi của sáng chế.

Theo các phương án của sáng chế, bộ máy có thể được phân chia thành các môđun chức năng dựa trên các ví dụ phương pháp trên. Ví dụ như, mỗi môđun chức năng có

thể thu được thông qua sự phân chia dựa trên mỗi chức năng tương ứng, hoặc hai hoặc nhiều chức năng có thể được tích hợp thành một môđun xử lý. Môđun tích hợp có thể được triển khai dưới dạng phần cứng, hoặc có thể được triển khai dưới dạng môđun chức năng phần mềm. Cần lưu ý rằng theo các phương án của sáng chế, việc phân chia thành các môđun là ví dụ và đơn thuần là sự phân chia chức năng logic, và có thể là sự phân chia khác trong triển khai thực tế. Ví dụ trong đó mỗi môđun chức năng thu được thông qua sự phân chia dựa trên mỗi chức năng tương ứng được sử dụng dưới đây cho sự mô tả.

Fig.14 là sơ đồ dạng giản đồ của cấu trúc của thực thể ML theo phương án của sáng chế. Như được thể hiện trên Fig.14, thực thể ML bao gồm đơn vị xử lý 101 và đơn vị truyền thông 102.

Một cách tùy chọn, thực thể ML có thể thực hiện bất kỳ một trong số các giải pháp sau.

Giải pháp 1

Thực thể ML hỗ trợ liên kết chính và ít nhất một liên kết không chính. Bộ đếm chờ truyền được bố trí trên liên kết chính, và không có bộ đếm chờ truyền được bố trí trên liên kết không chính. Đơn vị xử lý 101 được tạo cấu hình để thực hiện quy trình chờ truyền của liên kết chính dựa trên bộ đếm chờ truyền. Đơn vị truyền thông 102 được tạo cấu hình để: khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, được sử dụng bởi thực thể ML để gửi PPDU thứ nhất trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K liên kết thứ nhất bao gồm liên kết chính và K-1 liên kết không chính thứ nhất, liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0, và K là số nguyên dương.

Theo thiết kế khả thi, đơn vị truyền thông 102 được tạo cấu hình cụ thể để gửi PPDU thứ nhất trên kênh khả dụng của mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó kênh khả dụng của liên kết chính bao gồm kênh chính của liên kết chính, và kênh khả dụng của liên kết không chính thứ nhất bao gồm kênh chính của liên kết không chính thứ nhất.

Theo thiết kế khả thi, đơn vị xử lý 101 được tạo cấu hình cụ thể để: đợi chu kỳ nghỉ của kênh chính của liên kết chính để đạt tới khoảng thời gian liên khung thứ hai;

sau chu kỳ nghỉ của kênh chính của liên kết chính đạt tới khoảng thời gian liên khung thứ hai, mỗi lần kênh chính của liên kết chính ở trong trạng thái nghỉ trong một khe thời gian, giảm giá trị đếm của bộ đếm chờ truyền đi 1; và khi giá trị đếm của bộ đếm chờ truyền giảm xuống 0, kết thúc quy trình chờ truyền của liên kết chính.

Theo thiết kế khả thi, việc liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính bao gồm: Kênh chính của liên kết không chính thứ nhất ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm mà tại đó giá trị đếm của bộ đếm chờ truyền giảm xuống 0.

Theo thiết kế khả thi, kênh chính của liên kết không chính thứ nhất là kênh phụ có tần số thấp nhất, 20 MHz, trong dải tần số tương ứng với liên kết không chính thứ nhất. Ngoài ra, kênh chính của liên kết không chính thứ nhất là kênh phụ có tần số cao nhất, 20 MHz, trong dải tần số tương ứng với liên kết không chính thứ nhất.

Theo thiết kế khả thi, PPDU thứ nhất bao gồm khung MAC loại thứ nhất, trong đó khung MAC loại thứ nhất không cần đáp ứng.

Theo thiết kế khả thi, PPDU thứ nhất bao gồm khung MAC loại thứ hai, trong đó khung MAC loại thứ hai cần đáp ứng. Đơn vị truyền thông 102 còn được tạo cấu hình để nhận khung đáp ứng cho khung MAC loại thứ hai trên một hoặc nhiều liên kết thứ nhất. Đơn vị xử lý 101 còn được tạo cấu hình để: khi một hoặc nhiều liên kết thứ nhất không bao gồm liên kết chính, xác định rằng việc thiết lập TXOP thất bại; khi một hoặc nhiều liên kết thứ nhất bao gồm liên kết chính, xác định rằng việc thiết lập TXOP thành công.

Theo thiết kế khả thi, đơn vị xử lý 101 còn được tạo cấu hình để: xác định N liên kết thứ hai tương ứng với TXOP, trong đó N liên kết thứ hai bao gồm liên kết chính và N-1 liên kết không chính thứ hai, liên kết không chính thứ hai là liên kết không chính thứ nhất mà đáp ứng điều kiện đặt trước, trong đó điều kiện đặt trước bao gồm: Trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ nhất, hoặc trên liên kết không chính thứ nhất, thực thể ML gửi PPDU thứ nhất bao gồm khung MAC loại thứ hai, và nhận khung đáp ứng cho khung MAC loại thứ hai. Đơn vị truyền thông 102 còn được tạo cấu hình để gửi PPDU thứ hai trên mỗi liên kết thứ hai trong N liên kết thứ hai.

Theo thiết kế khả thi, đơn vị truyền thông 102 còn được tạo cấu hình để: nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết không chính thứ hai, dừng gửi PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thất bại, và tiếp tục gửi, đến khi TXOP kết thúc, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ hai thành công.

Theo thiết kế khả thi, đơn vị truyền thông 102 còn được tạo cấu hình để: nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, dừng gửi PPDU thứ hai trên N liên kết thứ hai. Đơn vị xử lý 101 còn được tạo cấu hình để đợi chu kỳ nghỉ của liên kết chính để đạt tới khoảng thời gian liên khung thứ nhất. Đơn vị truyền thông 102 còn được tạo cấu hình để: sau khi chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất, gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba, trong đó P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, liên kết không chính thứ ba là liên kết không chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất là thời điểm mà tại đó chu kỳ nghỉ của liên kết chính đạt tới khoảng thời gian liên khung thứ nhất, và P là số nguyên dương nhỏ hơn hoặc bằng N.

Theo thiết kế khả thi, đơn vị truyền thông 102 còn được tạo cấu hình để: nếu sự truyền PPDU thứ hai thất bại trên một hoặc nhiều liên kết thứ hai, dừng gửi PPDU thứ hai trên N liên kết thứ hai. Đơn vị xử lý 101 còn được tạo cấu hình để thực hiện quy trình chờ truyền của liên kết chính. Đơn vị truyền thông 102 còn được tạo cấu hình để: sau khi quy trình chờ truyền của liên kết chính kết thúc, gửi PPDU thứ hai trên mỗi liên kết thứ ba trong P liên kết thứ ba, trong đó P liên kết thứ ba bao gồm liên kết chính và P-1 liên kết không chính thứ ba, liên kết không chính thứ ba là liên kết không chính thứ hai mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết chính, và P là số nguyên dương nhỏ hơn hoặc bằng N.

Giải pháp 2

Thực thể ML hỗ trợ K liên kết thứ nhất. Đơn vị xử lý 101 được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Đơn vị truyền thông 102 được tạo cấu hình để: khi quy trình chờ truyền của liên kết mục tiêu kết thúc, gửi PPDU thứ nhất trên

mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, liên kết mục tiêu là liên kết thứ nhất mà quy trình chờ truyền kết thúc trước tiên trong K liên kết thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng K . Đơn vị truyền thông 102 còn được tạo cấu hình để: nếu sự truyền PPDU thứ nhất thất bại trên một hoặc nhiều liên kết thứ hai, bỏ qua gửi, ở trong chu kỳ thời gian đặt trước, PPDU thứ hai trên liên kết thứ hai mà trên đó sự truyền PPDU thứ nhất thất bại, hoặc bỏ qua gửi PPDU thứ hai trên N liên kết thứ hai ở trong chu kỳ thời gian đặt trước.

Giải pháp 3

Thực thể ML hỗ trợ K liên kết thứ nhất. Đơn vị xử lý 101 được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Đơn vị truyền thông 102 được tạo cấu hình để gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà quy trình chờ truyền đã kết thúc và ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất, và N là số nguyên dương nhỏ hơn hoặc bằng M .

Theo thiết kế khả thi, thời điểm thứ nhất là thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, và liên kết mục tiêu là liên kết thứ hai mà quy trình chờ truyền kết thúc cuối cùng trong N liên kết thứ hai.

Giải pháp 4

Thực thể ML hỗ trợ K liên kết thứ nhất. Đơn vị xử lý 101 được tạo cấu hình để thực hiện quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2. Đơn vị truyền thông 102 được tạo cấu hình để: khi tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất nhỏ hơn hoặc bằng 0, hoặc tổng các giá trị đếm của các bộ đếm chờ truyền của N liên kết thứ hai nhỏ hơn hoặc bằng 0, gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ hai trước thời điểm hiện tại, và N là số nguyên dương nhỏ hơn hoặc bằng M .

Theo thiết kế khả thi, đơn vị xử lý 101 được tạo cấu hình cụ thể để: đối với mỗi

liên kết thứ nhất trong K liên kết thứ nhất, đợi chu kỳ nghỉ của liên kết thứ nhất để đạt tới khoảng thời gian liên khung thứ hai; và sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, giảm giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất đi 1.

Theo thiết kế khả thi, giá trị đếm của bộ đếm chờ truyền của liên kết thứ nhất bao gồm số nguyên âm.

Theo thiết kế khả thi, đơn vị xử lý 101 còn được tạo cấu hình để: đối với mỗi liên kết thứ nhất trong K liên kết thứ nhất, sau khi chu kỳ nghỉ của liên kết thứ nhất đạt tới khoảng thời gian liên khung thứ hai, giảm giá trị đếm của bộ đếm mục tiêu đi 1 mỗi lần liên kết thứ nhất ở trong trạng thái nghỉ trong một khe thời gian, trong đó bộ đếm mục tiêu được tạo cấu hình để ghi lại tổng các giá trị đếm của các bộ đếm chờ truyền của K liên kết thứ nhất.

Giải pháp 5

Thực thể ML hỗ trợ nhiều liên kết, và nhiều liên kết mỗi liên kết đóng vai trò là liên kết thứ nhất đến lượt theo thứ tự tuần hoàn đặt trước. Đơn vị xử lý 101 được tạo cấu hình để thực hiện quy trình chờ truyền trên liên kết thứ nhất. Đơn vị truyền thông 102 được tạo cấu hình để: sau khi quy trình chờ truyền của liên kết thứ nhất kết thúc, gửi PPDU thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó N liên kết thứ hai bao gồm liên kết thứ nhất và N-1 liên kết khả dụng, liên kết khả dụng ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm kết thúc của quy trình chờ truyền của liên kết thứ nhất, và N là số nguyên dương.

Thực thể ML được đề xuất theo các phương án của sáng chế có thể được triển khai dưới nhiều dạng sản phẩm. Trong ví dụ, thực thể ML có thể được tạo cấu hình như là hệ thống xử lý chung. Trong ví dụ khác, thực thể ML có thể được triển khai sử dụng kiến trúc bus chung. Trong vẫn ví dụ khác, thực thể ML có thể được triển khai sử dụng mạch tích hợp chuyên dụng (application specific integrated circuit, ASIC). Phần sau đề xuất một vài dạng sản phẩm khả thi của thực thể ML theo các phương án của sáng chế. Cần hiểu rằng các dạng sản phẩm sau đơn thuần là các ví dụ, và các dạng sản phẩm khả thi của thực thể ML theo các phương án của sáng chế không bị giới hạn.

Fig.15 là sơ đồ kết quả của dạng sản phẩm khả thi của thực thể ML theo phương

án của sáng chế.

Theo dạng sản phẩm khả thi, thực thi ML theo các phương án của sáng chế có thể là thiết bị truyền thông, và thiết bị truyền thông bao gồm bộ xử lý 201 và bộ truyền nhận 202. Một cách tùy chọn, thiết bị truyền thông còn bao gồm phương tiện lưu trữ 203.

Bộ xử lý 201 được tạo cấu hình để thực hiện bước S101 trên Fig.5, bước S103 trên Fig.6, bước S201 trên Fig.9(a), bước S301 trên Fig.10, bước S401 trên Fig.11(a), bước S501 trên Fig.12, và bước S601 trên Fig.13. Bộ truyền nhận 202 được tạo cấu hình để thực hiện bước S102 trên Fig.5, bước S104 trên Fig.6, bước S202 và bước S203 trên Fig.9(a), bước S204 trên Fig.9(b), bước S302 trên Fig.10, bước S402 trên Fig.11(a), bước S403 trên Fig.11(b), bước S502 trên Fig.12, và bước S602 trên Fig.13.

Theo dạng sản phẩm khả thi khác, thực thể ML theo phương án này của sáng chế có thể còn được triển khai bởi bộ xử lý đa dụng hoặc bộ xử lý chuyên dụng mà được gọi thông thường là chip. Chip bao gồm mạch điện xử lý 201 và chân truyền nhận 202. Một cách tùy chọn, chip có thể còn bao gồm phương tiện lưu trữ 203.

Mạch điện xử lý 201 được tạo cấu hình để thực hiện bước S101 trên Fig.5, bước S103 trên Fig.6, bước S201 trên Fig.9(a), bước S301 trên Fig.10, bước S401 trên Fig.11(a), bước S501 trên Fig.12, và bước S601 trên Fig.13. Chân truyền nhận 202 được tạo cấu hình để thực hiện bước S102 trên Fig.5, bước S104 trên Fig.6, bước S202 và bước S203 trên Fig.9(a), bước S204 trên Fig.9(b), bước S302 trên Fig.10, bước S402 trên Fig.11(a), bước S403 trên Fig.11(b), bước S502 trên Fig.12, và bước S602 trên Fig.13.

Phương án của sáng chế còn đề xuất phương tiện lưu trữ đọc được bằng máy tính, trong đó phương tiện lưu trữ đọc được bằng máy tính lưu trữ các lệnh máy tính. Khi phương tiện lưu trữ đọc được bằng máy tính chạy trên thực thể ML, thực thể ML thực hiện phương pháp được thể hiện trên Fig.5, Fig.6, Fig.9(a), Fig.9(b), Fig.10, Fig.11(a), Fig.11(b), Fig.12, hoặc Fig.13. Các lệnh máy tính có thể được lưu trữ trong phương tiện lưu trữ đọc được bằng máy tính, hoặc có thể được truyền từ một phương tiện lưu trữ đọc được bằng máy tính đến phương tiện lưu trữ đọc được bằng máy tính khác. Ví dụ như, các lệnh máy tính có thể được truyền từ trang mạng, máy tính, máy chủ, hoặc trung tâm dữ liệu đến trang mạng khác, máy tính, máy chủ, hoặc trung tâm dữ liệu trong mạng có

dây (ví dụ như, cáp đồng trục, cáp quang, hoặc đường dây thuê bao kỹ thuật số (digital subscriber line, DSL)) hoặc không dây (ví dụ như, hồng ngoại, không dây hoặc vi ba). Phương tiện lưu trữ đọc được bằng máy tính có thể là bất kì phương tiện sử dụng được truy cập được bằng máy tính, hoặc thiết bị lưu trữ dữ liệu, như là máy chủ hoặc trung tâm dữ liệu, tích hợp một hoặc nhiều phương tiện sử dụng được. Phương tiện sử dụng được có thể là phương tiện từ (ví dụ như, đĩa mềm, đĩa cứng, hoặc băng từ), phương tiện quang (ví dụ, DVD), phương tiện bán dẫn (ví dụ như, ổ đĩa đặc (Solid State Disk, SSD)), hoặc tương tự.

Phương án của sáng chế còn đề xuất sản phẩm chương trình máy tính bao gồm các chỉ dẫn máy tính. Khi sản phẩm chương trình máy tính chạy trên thực thể ML, thực thể ML có thể thực hiện phương pháp được thể hiện trên Fig.5, Fig.6, Fig.9(a), Fig.9(b), Fig.10, Fig.11(a), Fig.11(b), Fig.12, hoặc Fig.13.

Mặc dù sáng chế được mô tả với tham chiếu đến các phương án, trong quá trình triển khai sáng chế mà yêu cầu bảo hộ, người có hiểu biết trung bình trong lĩnh vực có thể hiểu và triển khai các biến thể khác của các phương án được bộc lộ bằng việc xem các hình vẽ đi kèm, nội dung được bộc lộ, và các điểm bảo hộ đi kèm. Trong các điểm bảo hộ, thuật ngữ “bao gồm” không loại trừ thành phần khác hoặc bước khác, và “một” không loại trừ trường hợp số nhiều. Bộ xử lý đơn hoặc đơn vị khác có thể triển khai một vài chức năng được liệt kê trong các điểm bảo hộ. Một số sự đo lường được ghi lại trong các điểm phụ thuộc khác với nhau, nhưng không có nghĩa là các sự đo lường này không thể được kết hợp để tạo ra hiệu quả tốt hơn.

Mặc dù sáng chế được mô tả với sự tham chiếu đến các đặc tính cụ thể và các phương án của chúng, rõ ràng rằng các sửa đổi và các sự kết hợp khác nhau có thể được thực hiện cho chúng mà không xa rời khỏi phạm vi của sáng chế. Một cách tương ứng, bản mô tả và các hình vẽ đi kèm đơn thuần là mô tả ví dụ của sáng chế được quy định bởi các điểm bảo hộ đi kèm, và được coi như là bất kì hoặc tất cả các sửa đổi, biến thể, kết hợp, hoặc tương đương mà bao phủ phạm vi của sáng chế. Rõ ràng rằng người có hiểu biết trung bình trong lĩnh vực có thể thực hiện các sửa đổi và biến thể khác nhau cho sáng chế mà không xa rời khỏi phạm vi của sáng chế. Sáng chế nhằm để bao gồm các sửa đổi và biến thể của sáng chế với điều kiện là chúng nằm trong phạm vi bảo hộ được quy định bởi các điểm yêu cầu bảo hộ của sáng chế và công nghệ tương đương của

chúng.

YÊU CẦU BẢO HỘ

1. Phương pháp truyền thông, trong đó phương pháp được áp dụng cho thực thể đa liên kết (multi-link, ML), thực thể ML hỗ trợ K liên kết thứ nhất, và phương pháp còn bao gồm:

bước thực hiện (S301), bởi thực thể ML, quy trình chờ truyền trên mỗi liên kết thứ nhất trong K liên kết thứ nhất, trong đó K là số nguyên dương lớn hơn hoặc bằng 2; và

bước gửi (S302), bởi thực thể ML, đơn vị dữ liệu giao thức lớp vật lý (physical layer protocol data unit, PPDU) thứ nhất trên mỗi liên kết thứ hai trong N liên kết thứ hai, trong đó liên kết thứ hai là liên kết thứ nhất mà quy trình chờ truyền đã kết thúc và ở trong trạng thái nghỉ trong khoảng thời gian liên khung thứ nhất trước thời điểm thứ nhất, và N là số nguyên dương lớn hơn hoặc bằng 2 và nhỏ hơn hoặc bằng K; và

thời điểm thứ nhất là thời điểm kết thúc của quy trình chờ truyền của liên kết mục tiêu, và liên kết mục tiêu là liên kết thứ hai mà quy trình chờ truyền kết thúc cuối cùng trong N liên kết thứ hai.

2. Phương pháp truyền thông theo điểm 1, bước gửi (S302), bởi thực thể ML, đơn vị dữ liệu giao thức lớp vật lý thứ nhất, PPDU trên mỗi liên kết thứ hai trong N liên kết thứ hai, bao gồm:

bước gửi (S302) đồng thời, bởi thực thể ML, PPDU trên mỗi liên kết thứ hai trong N liên kết thứ hai.

3. Thực thể đa liên kết (multi-link, ML), trong đó thực thể ML bao gồm đơn vị được tạo cấu hình để thực hiện các bước trong phương pháp theo điểm 1 hoặc 2.

4. Thực thể đa liên kết (multi-link, ML) theo điểm 3, trong đó thực thể ML là thực thể trạm (station, STA) đa liên kết hoặc trong đó thực thể ML bao gồm nhiều STA.

5. Phương tiện lưu trữ đọc được bằng máy tính, trong đó phương tiện lưu trữ đọc được bằng máy tính bao gồm các lệnh máy tính; và khi các lệnh máy tính được chạy trên máy tính, máy tính thực hiện phương pháp truyền thông theo điểm 1 hoặc 2.

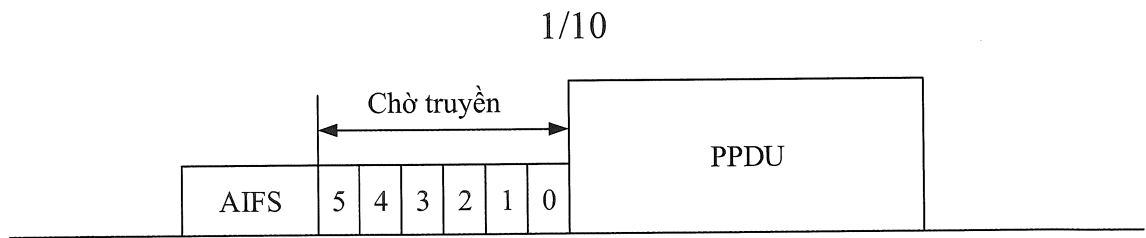


FIG. 1

2/10

L-STF	L-LTF	L-SIG	RL-SIG	HE-SIG-A	HE-STF	HE-LTF	...	HE-LTF	DATA	PE
-------	-------	-------	--------	----------	--------	--------	-----	--------	------	----

FIG. 2

3/10

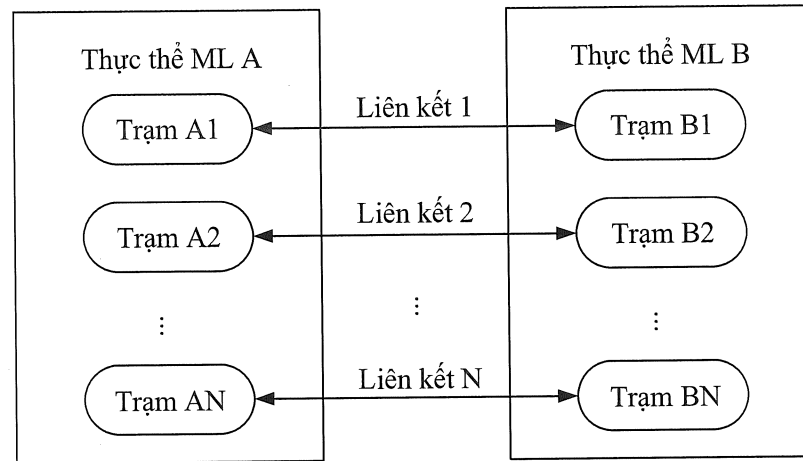


FIG. 3

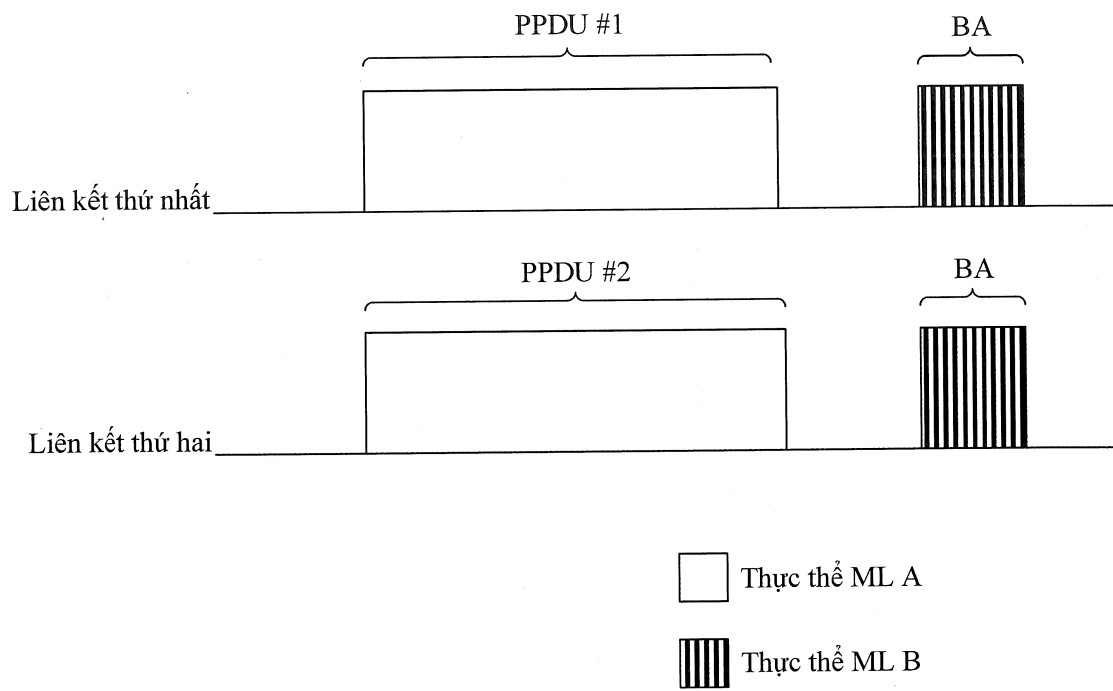


FIG. 4

4/10

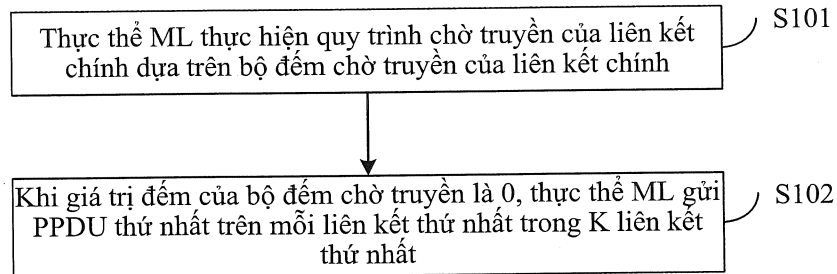


FIG. 5

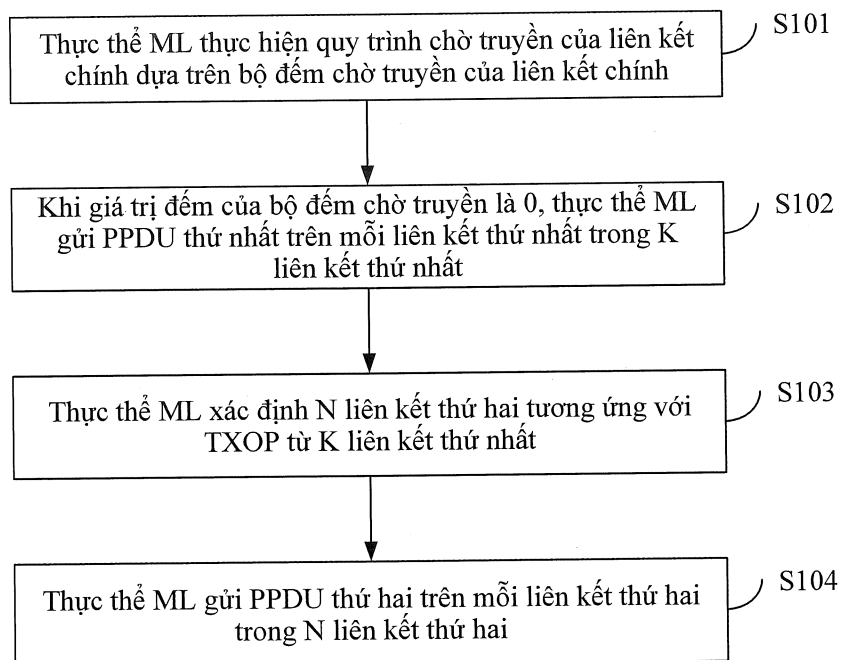


FIG. 6

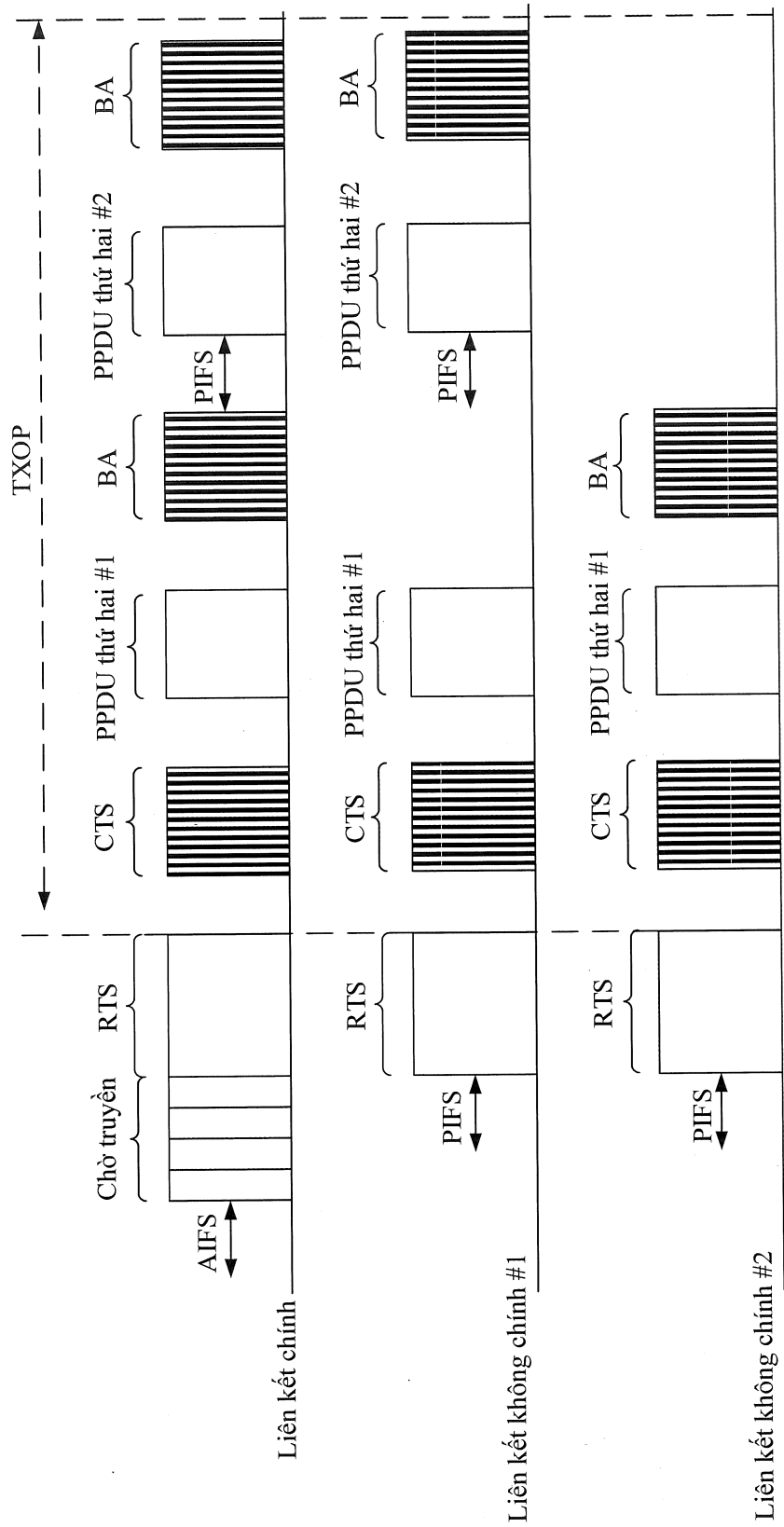


FIG. 7

6/10

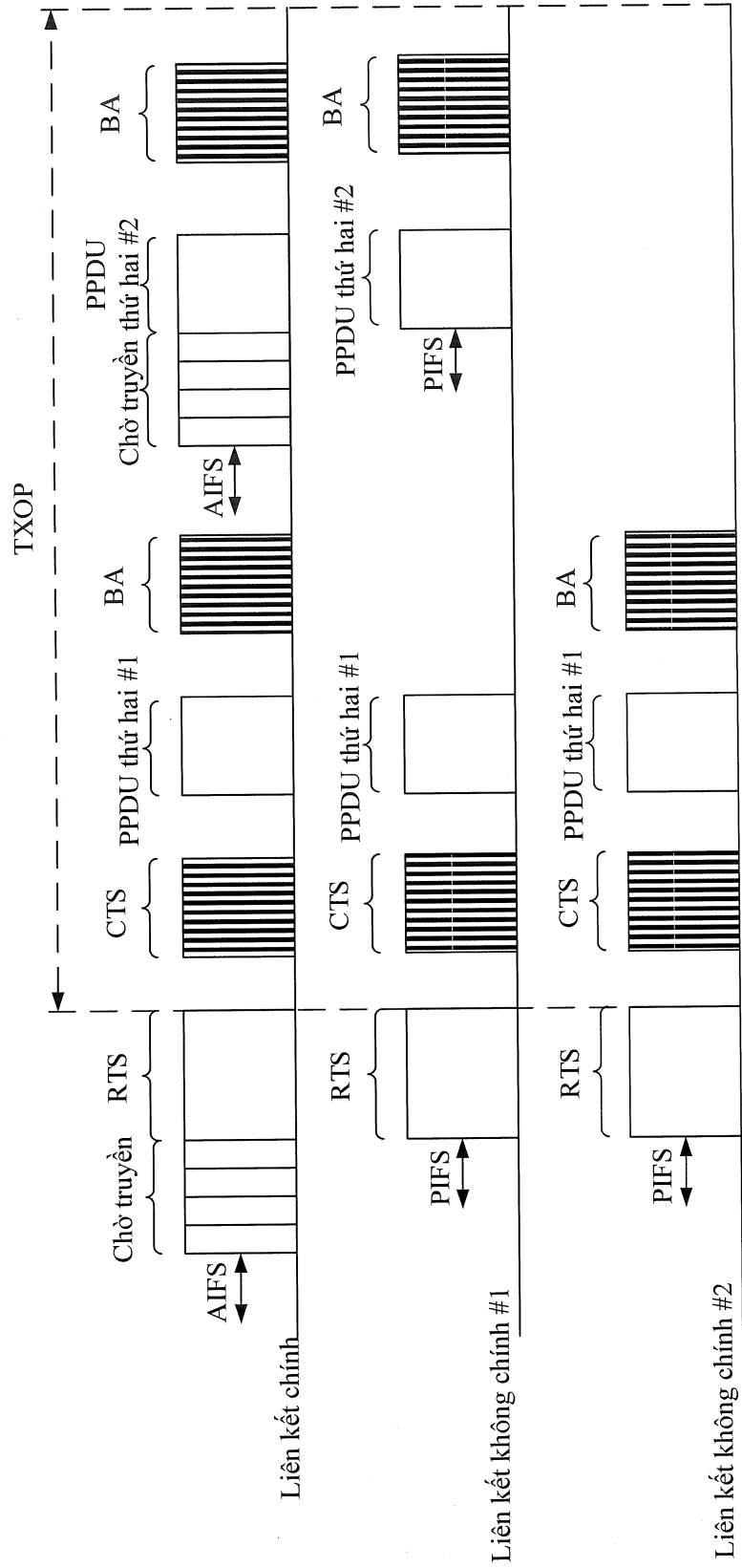


FIG. 8

7/10

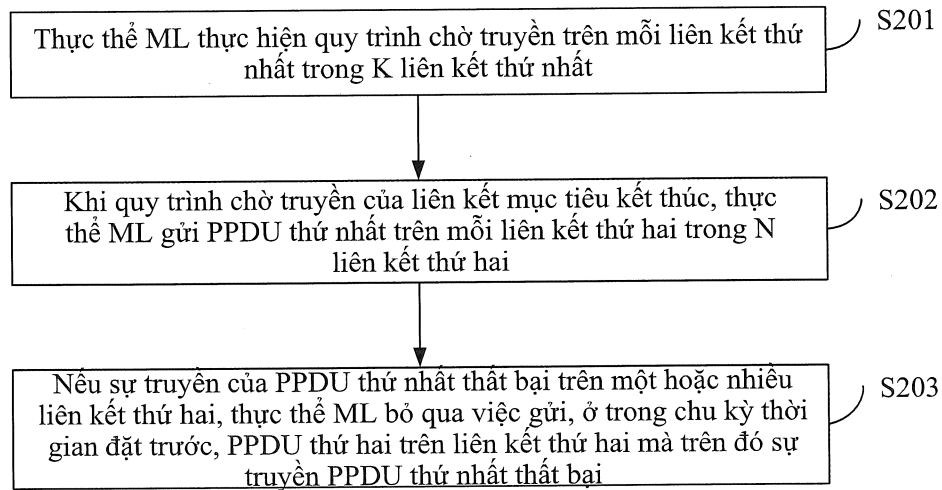


FIG. 9(a)

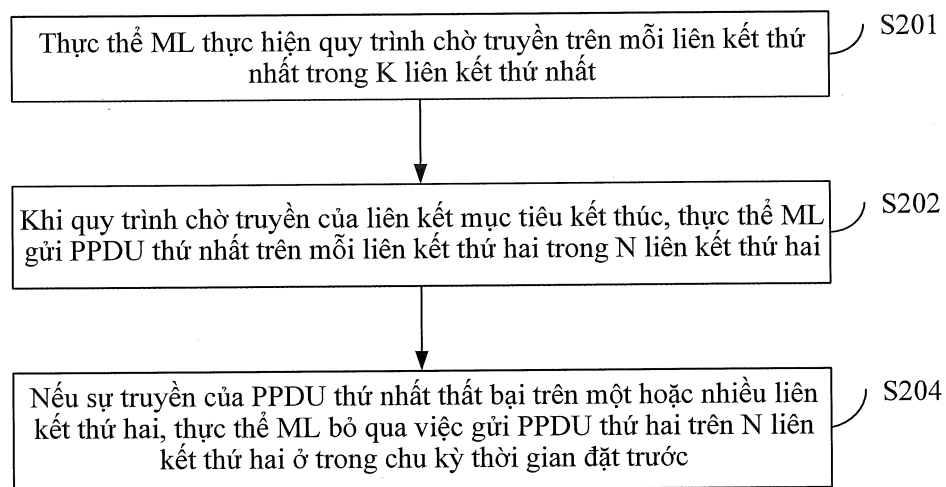


FIG. 9(b)

8/10

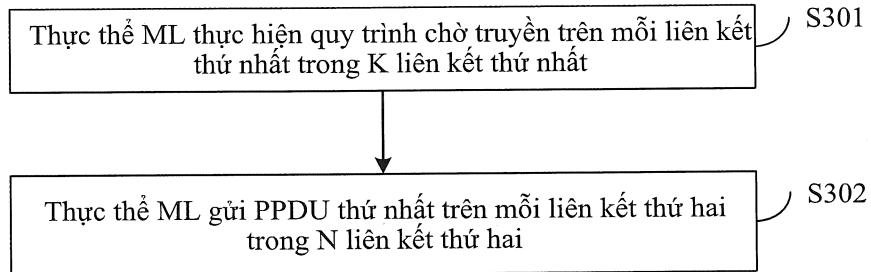


FIG. 10

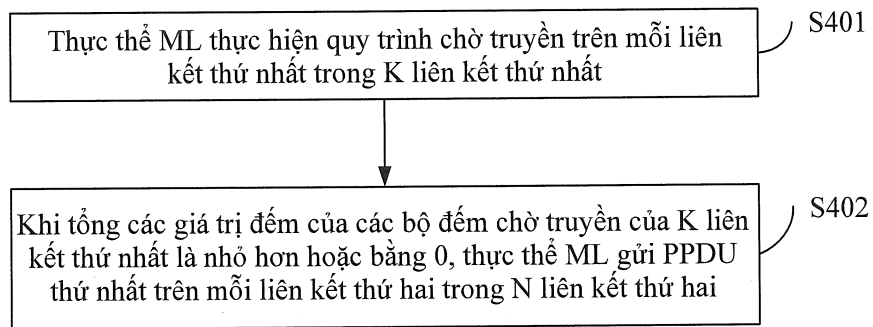


FIG. 11(a)

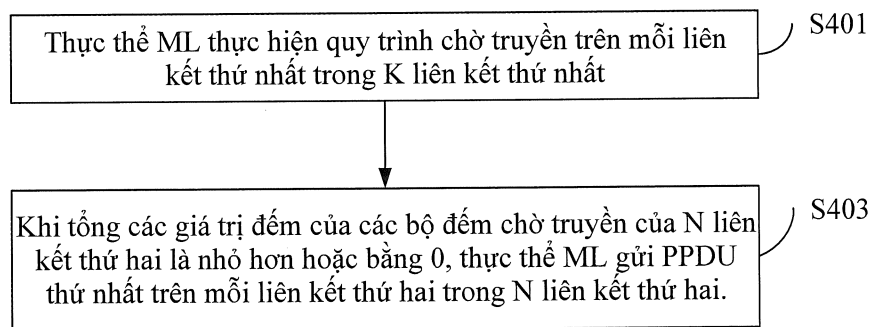


FIG. 11(b)

9/10

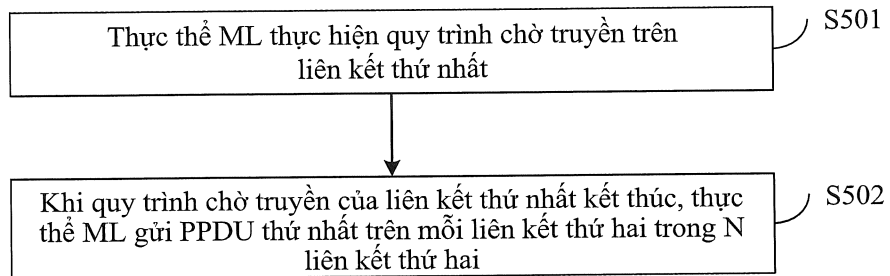


FIG. 12

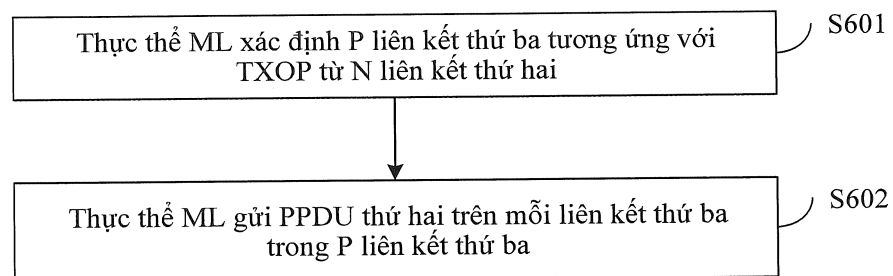


FIG. 13

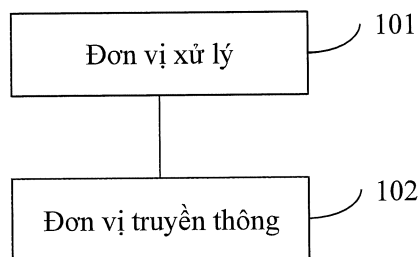


FIG. 14

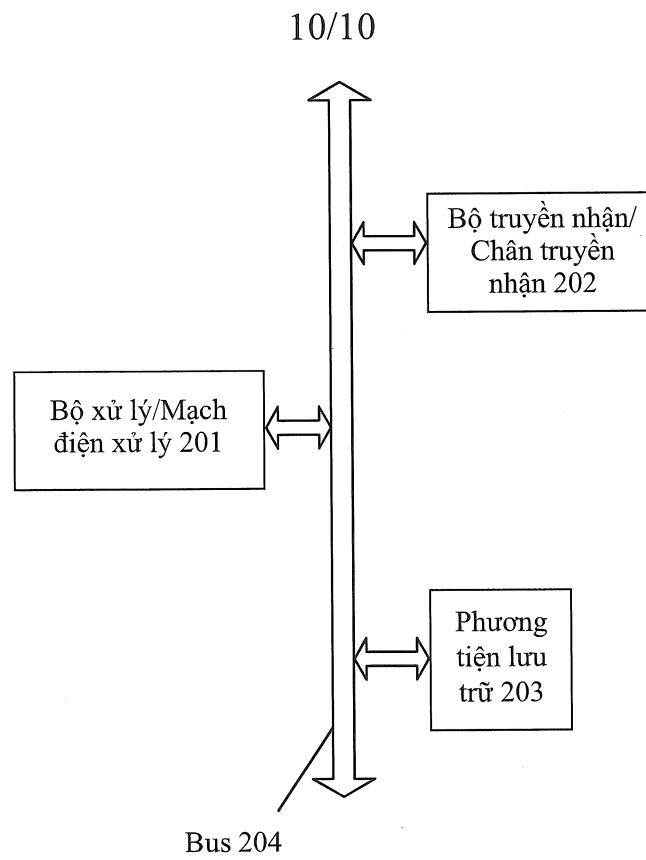


FIG. 15